

Log-Concave Sampling

Sinho Chewi

This Version: August 9, 2026

Preface	vi
Part I Diffusions in Continuous Time	1
1 The Langevin Diffusion in Continuous Time	3
1.1 A Primer on Stochastic Calculus	3
1.2 Markov Semigroup Theory	10
1.3 The Geometry of Optimal Transport	19
1.4 The Langevin SDE as a Wasserstein Gradient Flow	31
1.5 Overview of the Convergence Results	36
Bibliographical Notes	40
Exercises	41
2 Functional Inequalities	48
2.1 Overview of the Inequalities	48
2.2 Proofs via Markov Semigroup Theory	50
2.3 Operations Preserving Functional Inequalities	60
2.4 Concentration of Measure and Isoperimetry	68
2.5 Riemannian Manifolds	77
2.6 Discrete Space and Time	84
Bibliographical Notes	85
Exercises	87
3 Additional Topics in Stochastic Analysis	96
3.1 Quadratic Variation	96
3.2 Change of Measure in Path Space	100
3.3 Doob's Transform	104
3.4 Föllmer Drift	108
3.5 Schrödinger Bridge	110
Bibliographical Notes	114

Exercises	116
Part II Complexity of Sampling	121
4 Analysis of Langevin Monte Carlo	123
4.1 Proof via Wasserstein Coupling	124
4.2 Proof via Interpolation Argument	127
4.3 Proof via Convex Optimization	131
4.4 Proof via Girsanov's Theorem	134
Bibliographical Notes	137
Exercises	138
5 Faster Low-Accuracy Samplers	141
5.1 Randomized Midpoint Discretization	141
5.2 Hamiltonian Monte Carlo	145
5.3 The Underdamped Langevin Diffusion	149
Bibliographical Notes	166
Exercises	168
6 Convergence in Rényi Divergence	174
6.1 Analysis of LMC via Interpolation Argument	174
6.2 Analysis of LMC via Girsanov's Theorem	178
6.3 Analysis of ULMC via Girsanov's Theorem	183
Bibliographical Notes	184
Exercises	185
7 High-Accuracy Samplers	189
7.1 Rejection Sampling	189
7.2 The Metropolis–Hastings Filter	191
7.3 An Overview of High-Accuracy Samplers	193
7.4 Markov Chains in Discrete Time	197
7.5 Analysis of MALA for a Cold Start	203
7.6 Analysis of MALA for a Warm Start	206
Bibliographical Notes	209
Exercises	211
8 The Proximal Sampler	214
8.1 Introduction to the Proximal Sampler	215
8.2 Convergence under Strong Log-Concavity	216
8.3 Simultaneous Flow and Time Reversal	218
8.4 Convergence under Log-Concavity	221
8.5 Convergence under Functional Inequalities	222
8.6 Implementations of the RGO and Applications	223
Bibliographical Notes	229
Exercises	230
9 Lower Bounds for Sampling	233
9.1 A Brief Discussion of Query Complexity	233

9.2	Query Complexity in One Dimension	234
9.3	Query Complexity in Constant Dimension	241
9.4	Query Complexity for Gaussians	243
	Bibliographical Notes	245
	Exercises	246
10	Structured Sampling	249
10.1	Stochastic Gradients	249
10.2	Coordinate Methods	252
10.3	Mirror Langevin	257
	Bibliographical Notes	268
	Exercises	269
11	Non-Log-Concave Sampling	272
11.1	Approximate First-Order Stationarity via Fisher Information	272
11.2	Fisher Information Bounds	274
11.3	Applications of Fisher Information Bounds	276
11.4	Lower Bounds	279
	Bibliographical Notes	280
	Exercises	280
12	Diffusion Generative Models	283
12.1	Introduction	283
12.2	Score Matching and Variants	284
12.3	Discretization Analysis	288
	Bibliographical Notes	295
	References	297
	<i>Index</i>	315

What This Book Contains

As discussed in the next section, a large portion of this book is dedicated to a systematic and unified treatment of recent developments in the complexity theory for log-concave sampling, with a particular emphasis on connections with the field of optimization. Many of these developments appear here in textbook form for the first time. Although this is still an active area of research, at this time there is enough elegant mathematics and canonical theory that it seemed a shame not to have an exposition accessible to an ambitious graduate student.

From a broader view, however, it is not the specific applications to log-concave sampling, but rather the general perspective and techniques used, that will have the largest impact on the reader. With this in mind, the book includes several topics which are not directly related to sampling, but loosely illustrate the general theme of “modern applications of stochastic analysis to probability and statistics”. The topics range from classical mathematical questions in concentration of measure and geometry to recent machine learning algorithms inspired by diffusions. Although tastes change, the overall importance of this perspective only seems to grow with time. Due to space limitations, some of these additional topics appear as supplements on my website (chewisinho.github.io).

The subject matter of this book touches upon many fields, including functional analysis, geometry, partial differential equations, and stochastic calculus, and a primary goal of the exposition here is to make the material accessible without extensive background knowledge in these topics. This means that in several places we have sacrificed full mathematical rigor in favor of (hopefully) more lucid explanations, referring to the original sources for details. As such, these subjects are not prerequisites for this book, although more background knowledge on the reader’s part naturally translates into a healthier understanding of the context of the material. The main exceptions to this statement are: (1) we assume that the reader is familiar with graduate-level analysis and probability; (2) since much of the theory of sampling is inspired by ideas from optimization, we strongly recommend that the reader is familiar with the latter, as treated in, e.g., [Bubeck \(2015\)](#); [Nesterov \(2018\)](#); [Chewi \(2025\)](#).

The Complexity of Sampling

In this book, we consider the following canonical sampling problem:

Given query access to a smooth function $V : \mathbb{R}^d \rightarrow \mathbb{R}$, what is the minimum number of queries required to output an approximate sample from the probability density $\pi \propto \exp(-V)$ on $\mathbb{R}^d$?

The problem formulation is chosen due to the following considerations. In many applications (some described below), we wish to sample from a probability density π , and we have an explicit function $V : \mathbb{R}^d \rightarrow \mathbb{R}$ such that $\pi \propto \exp(-V)$. In other words, since π is a probability density, then $\pi = \frac{1}{Z} \exp(-V)$, where $Z := \int \exp(-V)$ is called the *normalizing factor* (or the *partition function* in statistical physics). Although the function V specifies the density π up to normalization, a naïve evaluation of Z as a high-dimensional integral is intractable. Indeed, the usual approach of approximating an integral by a sum requires discretizing space via a fine grid whose size scales exponentially in the dimension d . Moreover, even if we had access to Z , it is still not clear how we could use this to sample from π . Therefore, the focus here is to develop direct methods for the sampling task which bypass the computation of Z .¹

Not only do we want to develop fast algorithms, we also want to understand the inherent complexity of the sampling task, which in turn allows us to identify *optimal* algorithms. By *complexity*, we do not mean unconditional lower bounds in *computational* complexity, which would quickly run into issues such as bit representations. Instead, following the well-trodden path of optimization, we adopt a model in which we only have access to V through queries made to an oracle, and our notion of complexity is the number of queries made. This is known as *oracle complexity* or *query complexity*; see Nemirovsky and Yudin (1983, §1) for a detailed discussion. We usually consider a first-order oracle, i.e., given a point $x \in \mathbb{R}^d$, the oracle returns $(V(x), \nabla V(x))$. Since V is only well-defined up to an additive constant, we can equivalently imagine that the oracle returns $(V(x) - V(0), \nabla V(x))$.

Before considering the problem further, here are some important applications.

- 1 **(Bayesian statistics and uncertainty quantification.)** Suppose that we wish to make inferences about a parameter θ of interest, which lies in a space $\Theta \subseteq \mathbb{R}^d$. As a Bayesian, we have a *prior* density p_θ over Θ which encodes our subjective beliefs about the value of θ prior to seeing any data. Next, we collect some data X which, conditionally on the value of θ , is drawn from a density $p_{X|\theta}(\cdot | \theta)$. According to Bayesian statistics, we should then compute the *posterior distribution*

$$p_{\theta|X}(\theta | X) = \frac{p_\theta(\theta) p_{X|\theta}(X | \theta)}{\int_{\Theta} p_\theta(d\theta') p_{X|\theta'}(X | \theta')},$$

which encodes our new beliefs about θ after seeing the data.

Typically, we have access to the functional forms of p_θ and $\{p_{X|\theta}(\cdot | \theta)\}_{\theta \in \Theta}$, so we can often evaluate these densities and compute gradients. However, the denominator of $p_{\theta|X}$ is precisely the normalizing constant described previously and cannot be easily evaluated. Moreover, even if we had the functional form of $p_{\theta|X}(\cdot | X)$, we would still not be able to compute expectations $\mathbb{E}[\varphi(\theta) | X]$ of test functions w.r.t. the posterior without evaluating another high-dimensional integral. Instead, the sampling methods we discuss in this book can output random variables $\vartheta_1, \dots, \vartheta_m$ approximately distributed according to $p_{\theta|X}(\cdot | X)$, and the expectation can be approximated to arbitrary accuracy via the Monte Carlo averages $m^{-1} \sum_{i=1}^m \varphi(\vartheta_i)$.

- 2 **(High-dimensional integration.)** More generally, computing integrals of functions against a known density π is a fundamental task in scientific computing. In many high-dimensional applications, the strategy of drawing samples from π and then approximating integrals via Monte Carlo averages is

¹ In fact, it goes the other way around: many state-of-the-art methods for approximately computing Z are based on sampling methods (Dyer et al., 1991; Cousins and Vempala, 2018).

one of the few general-purpose methods whose computational cost does not deteriorate exponentially with the dimension.

- 3 **(Model aggregation, online learning, etc.)** Consider the following regression model: we observe i.i.d. pairs $\{(X_i, Y_i)\}_{i \in [n]}$, assumed to be drawn according to $Y_i = f_\theta(X_i) + \xi_i$ for some parameter $\theta \in \Theta$ and noise variables $\xi_i, i \in [n]$. Given a “prior” p_θ over Θ , the *exponentially weighted aggregate* (EWA) estimator with parameter $\beta > 0$ is the following weighted average:

$$\hat{f} := \int f_\theta \hat{p}_{\beta, \theta}(\mathrm{d}\theta), \quad \text{where} \quad \hat{p}_{\beta, \theta}(\mathrm{d}\theta) \propto \exp\left(-\beta \sum_{i=1}^n (Y_i - f_\theta(X_i))^2\right) p_\theta(\mathrm{d}\theta).$$

Note the similarity with the Bayesian setting; indeed, $\hat{p}_{\beta, \theta}$ can be interpreted as a posterior. The estimator $\hat{f}$, however, can be viewed through a purely frequentist lens, as an average over “models” $\theta \in \Theta$ which downweights those which incur large empirical risk. This procedure is known generally as *model aggregation*, and the risk bound developed for the EWA estimator in Dalalyan (2017b) is closely related to *PAC–Bayes theory* (Alquier, 2024). Similar weighted averages appear in the field of online learning, under various names such as *learning with expert advice* or *multiplicative weights*. Computation of these estimators can again be handled via sampling methods.

- 4 **(Privacy.)** As machine learning algorithms are continually deployed in application domains involving personal and sensitive information, there is growing concern about maintaining the privacy of the data on which the machine learning models are trained. One way to address this issue is to require that the algorithm be *differentially private* (Dwork and Roth, 2013), which requires that the law of the output of the model does not depend too much on the presence or absence of a single data point. The most common method to achieve this goal is via the careful addition of noise to the algorithm. A precise connection between the fields is the exponential mechanism (McSherry and Talwar, 2007), which ensures privacy by sampling from a certain distribution. Readers who are interested in the mathematics of privacy should benefit from a healthy understanding of the analysis of sampling algorithms, and vice versa.
- 5 **(Statistical physics.)** In a physical system, $V(x)$ represents the energy of a state x . In this situation, thermodynamics predicts that the equilibrium distribution over states is the *Boltzmann* (or *Gibbs*) distribution whose density is proportional to $\exp(-V/T)$ (where T is the temperature of the system). Naturally, sampling provides a method for probing properties of the equilibrium distribution. More subtly, the mixing time of specific sampling algorithms also provides information about the system such as metastability phenomena; we revisit this in Chapter 11.

Due to this physical interpretation, we often refer to V as the *potential energy*.

Besides these examples, it is no surprise that sampling arises in many other applications, since sampling is a fundamental algorithmic primitive. As such, sampling methods are widely used in applied domains such as biology, climatology, and cosmology.

Example 0.0.1 (Bayesian logistic regression). For concreteness, let us consider the application of sampling to a Bayesian logistic regression problem. Suppose we have collected data in the form of pairs (X_i, Y_i) , $i \in [n]$, where $X_i \in \mathbb{R}^d$ is a vector of covariates and $Y_i \in \{0, 1\}$ is a binary outcome. For example, Y_i might represent whether or not a certain drug is effective on the i -th patient in a clinical study. Here, we regard the covariates $\{X_i\}_{i \in [n]}$ to be deterministic and fixed, and we posit that the outcomes $\{Y_i\}_{i \in [n]}$ are independent with distributions

$$Y_i \sim \text{Bernoulli}\left(\frac{\exp \langle \vartheta, X_i \rangle}{1 + \exp \langle \vartheta, X_i \rangle}\right).$$

Moreover, we take a Gaussian prior $\text{normal}(0, \lambda^{-1} I_d)$ for ϑ , where $\lambda > 0$. The likelihood is

$$p_{Y_i|\vartheta}(y_i | \vartheta) = \left(\frac{1}{1 + \exp \langle \vartheta, X_i \rangle}\right)^{1-y_i} \left(\frac{\exp \langle \vartheta, X_i \rangle}{1 + \exp \langle \vartheta, X_i \rangle}\right)^{y_i}, \quad y_i \in \{0, 1\},$$

and by independence, $p_{Y_1, \dots, Y_n|\vartheta}(y_1, \dots, y_n | \vartheta) = \prod_{i=1}^n p_{Y_i|\vartheta}(y_i | \vartheta)$. A Bayes rule computation yields the posterior $p_{\vartheta|Y_1, \dots, Y_n} \propto \exp(-V)$, where

$$V(\vartheta) = \sum_{i=1}^n (\log(1 + \exp \langle \vartheta, X_i \rangle) - Y_i \langle \vartheta, X_i \rangle) + \frac{\lambda}{2} \|\vartheta\|^2.$$

Note that V is λ -strongly convex and $(\lambda + \frac{1}{4} \|\sum_{i=1}^n X_i X_i^\top\|_{\text{op}})$ -smooth. It is straightforward to find the minimizer of V via basic optimization methods (e.g., gradient descent), and this corresponds to finding the mode or *maximum a posteriori* (MAP) estimate of the parameter ϑ . On the other hand, it is less obvious how to obtain (approximate) samples from the posterior. In this book, we study algorithms that solve this task, together with non-asymptotic complexity estimates.

Next, we turn toward the *how* rather than the *why*. A key theme of this book is the surprising and close connection between methods in *optimization* and methods in *sampling*. To illustrate, we introduce our first sampling method, which is the sampling analogue of the well-known gradient descent algorithm from optimization. The **Langevin diffusion** $(X_t)_{t \geq 0}$ is defined by the stochastic differential equation (SDE)

$$dX_t = \underbrace{-\nabla V(X_t) dt}_{\text{gradient flow}} + \underbrace{\sqrt{2} dB_t}_{\text{Brownian motion}}.$$

With a pure gradient flow $dX_t = -\nabla V(X_t) dt$, we would expect the dynamics to converge to stationary points of V . The Brownian motion ensures that we fully explore the distribution π , as is required in sampling. Under mild conditions, the unique stationary distribution of the Langevin diffusion is indeed $\pi \propto \exp(-V)$, which makes this diffusion a good candidate upon which to base a sampling algorithm.

Since the Langevin diffusion is “a gradient flow + noise”, it is no wonder that researchers have drawn parallels between this diffusion and the gradient flow from optimization. However, the connection actually lies much deeper than this superficial observation would suggest. There is a natural geometry on the space of probability measures with finite second moment, $\mathcal{P}_2(\mathbb{R}^d)$, namely the 2-Wasserstein distance W_2 from the theory of optimal transport. The space $(\mathcal{P}_2(\mathbb{R}^d), W_2)$ turns out to be much richer than a metric space; in fact, it is *almost* a Riemannian manifold (Otto, 2001; Ambrosio et al., 2008). In turn, the Riemannian structure allows us to define *gradient flows* on this space. The punchline here is that if π_t denotes the law of X_t , then the curve of measures $t \mapsto \pi_t$ is the gradient flow of the Kullback–Leibler (KL) divergence $\text{KL}(\cdot \| \pi)$ with respect to the W_2 geometry. Hence, at the level of

the trajectory $(X_t)_{t \geq 0}$, the Langevin diffusion is a noisy gradient flow, but at the level of measures $(\pi_t)_{t \geq 0}$, it is *precisely* a gradient flow! This remarkable connection was introduced in the seminal work of Jordan, Kinderlehrer, and Otto (Jordan et al., 1998).

This perspective suggests that we can study the convergence of the Langevin diffusion using tools from optimization. For example, a standard assumption in the optimization literature which allows for fast rates of convergence is that of *strong convexity* of the objective function. Hence, we can ask under what conditions the functional $\text{KL}(\cdot \parallel \pi)$ on the space of measures is strongly convex along W_2 geodesics. Quite pleasingly, this is equivalent to the (Euclidean) strong convexity of the potential V . Consequently, the assumption of strong convexity of V , which is natural in optimization for studying gradient flows, turns out to be natural in the sampling context as well.

Much of this book is devoted to the case when V is strongly convex; we refer to π as being *strongly log-concave*. Besides its naturality and simplicity, it is sometimes a practical assumption. For example, in the application to Bayesian statistics, the Bernstein–von Mises theorem states that as the number of data points tends to infinity, the posterior distribution is asymptotically Gaussian, and is therefore expected to be (locally) strongly log-concave. This has already been exploited to give sampling guarantees in the context of Bayesian inverse problems (see, e.g., Nickl and Wang, 2024). However, some of the results also apply to restricted classes of non-log-concave measures, and in Chapter 11 we shall see what can be said about non-log-concave sampling in general.

To use the Langevin diffusion as an algorithm, however, we must first discretize the process in time. The simplest discretization, known as the *Euler–Maruyama discretization*, proceeds by fixing a step size $h > 0$ and following the iteration

$$\hat{X}_{(n+1)h} := \hat{X}_{nh} - h \nabla V(\hat{X}_{nh}) + \sqrt{2} (B_{(n+1)h} - B_{nh}).$$

Since the Brownian increment $B_{(n+1)h} - B_{nh}$ has the $\text{normal}(0, hI_d)$ distribution, this iteration can be easily implemented once we have access to a gradient oracle for V and the ability to draw standard Gaussian variables. This iteration is commonly known as the *Langevin Monte Carlo* (LMC) algorithm, or the *unadjusted Langevin algorithm* (ULA); in this book, we stick to the former acronym.

The LMC algorithm is the starting point of our study. As a result of research in the last decade, we now have the following guarantee. For any of the common divergences d between probability measures, e.g., $d(\mu, \pi) = W_2(\mu, \pi)$ or $d(\mu, \pi) = \sqrt{\text{KL}(\mu \parallel \pi)}$, and with an appropriate choice of initialization and step size, the law $\hat{\mu}_{Nh}$ of the N -th iterate (or mixture of iterates) of LMC satisfies $d(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$ with a number of iterations N which is polynomial in the problem parameters (the dimension d , the condition number κ of V , and the inverse accuracy $1/\varepsilon$). For example, when $d = \sqrt{\text{KL}}$, the state-of-the-art guarantee reads $N = \tilde{O}(\kappa d / \varepsilon^2)$ (Theorem 4.3.6). Whereas the convergence of the continuous-time diffusion is classical and typically proven via abstract calculus, the quantitative non-asymptotic convergence of the discretized algorithm necessitates the development of a new toolbox of analysis techniques. A primary goal of this book is to make this toolbox more accessible to researchers who are not yet acquainted with the field.

Beyond the standard LMC algorithm, there is now a rich arsenal of algorithms in the sampling literature. Some algorithms are directly inspired by other optimization algorithms (e.g., mirror descent), whereas other algorithms have their roots in the classical theory of Markov processes (e.g., the use of a Metropolis–Hastings filter). We also explore some of these more sophisticated algorithms in detail, as in many cases they represent substantial improvements over standard LMC.

Finally, although we began this preface by discussing the goal of understanding the *complexity* of sampling, in fact the complexity is not yet fully understood. The issue here is that there are currently

very few *lower bounds* on the complexity of sampling. This is in contrast with the field of optimization, in which oracle complexity lower bounds have in most situations identified nearly optimal algorithms for optimizing various function classes. In Chapter 9, we will explain the current progress toward achieving this goal for sampling, but much work remains to be done.

For example, here is the precise statement for a fundamental open question about the complexity of log-concave sampling.

Let $\pi \propto \exp(-V)$ be a probability density on $\mathbb{R}^d$. Determine, up to a universal constant, the minimum number of queries to a first-order oracle for V required to output a sample whose law μ satisfies $\|\mu - \pi\|_{\text{TV}} \leq \varepsilon$, uniformly over the following class of potentials: V is twice continuously differentiable, satisfying the conditions $0 < I_d \leq \nabla^2 V \leq \kappa I_d$ and $\nabla V(0) = 0$.

What This Book Does Not Contain

Here, we point out a few egregious exclusions from our selection of topics. First, as mentioned previously, the price we paid for a succinct exposition of a variety of fields is a lack of rigorous development of the fundamentals of said fields, which we leave to the reader to pursue more thoroughly.

The field of sampling has a rich literature spanning decades, and although we have made an effort to cite the works most relevant to the modern perspective, it was not possible to cite even a vanishing fraction of the applied and/or classical literature. This extends even to recent theoretical works on log-concave sampling, for which we have omitted discussion of sampling from convex bodies or polytopes (as initiated in [Dyer et al., 1991](#)). Although these works constitute fundamental developments in the field, here we focus on the strand of literature which is most strongly inspired by optimization algorithms.

Naturally, the other topics we explore in the book are far from comprehensive.

Notational Conventions

The symbols $\wedge$ and $\vee$ mean “minimum” and “maximum”, respectively. We write $a \lesssim b$ or $a = O(b)$ to mean that $a \leq Cb$ for a universal constant $C > 0$. Similarly, $a \gtrsim b$ or $a = \Omega(b)$ mean that $a \geq cb$ for a universal constant $c > 0$, and $a \asymp b$ or $a = \Theta(b)$ mean that both $a \lesssim b$ and $a \gtrsim b$. We write $a = \tilde{O}(b)$ to mean that $a = O(b \log^{O(1)} b)$, i.e., we suppress polylogarithmic factors, and we similarly use the notation $\tilde{\Omega}$ and $\tilde{\Theta}$. If a theorem holds under the condition “ $a \ll b$ ”, it means that it holds provided that a is smaller than b by a sufficiently small universal constant, and “ $a \ll_{\log} b$ ” has a similar meaning except that polylogarithmic factors are also ignored.

For a function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, we write $\partial_i f$ to denote the i -th partial derivative of f . The gradient ∇f is the vector of partial derivatives $(\partial_1 f, \dots, \partial_d f)$, and the Hessian $\nabla^2 f$ is the matrix $(\partial_i \partial_j f)_{i,j \in [d]}$. For a vector field $v : \mathbb{R}^d \rightarrow \mathbb{R}^d$, we also use the notation ∇v to denote the Jacobian matrix of v . The divergence of a vector field v is $\text{div } v = \nabla \cdot v = \sum_{i=1}^d \partial_i v_i$, and the Laplacian of f is $\Delta f = \text{tr } \nabla^2 f = \sum_{i=1}^d \partial_i^2 f$.

Finally, we sometimes use t_- to denote $\lfloor t/h \rfloor h$, the largest multiple of the step size h which is at most t . This is not to be confused with the negative part of t .

Acknowledgements

This book started off as a series of lecture notes during my visit to New York University in Spring 2022, and owes much to the audience and hospitality there. The notes were further developed during my time as a member at the Institute for Advanced Study, from 2023 through 2024, during which I was supported by the Eric and Wendy Schmidt Fund.

The idea for the book was first conceived during the Simons Institute for the Theory of Computing program on Geometric Methods in Sampling and Optimization (GMOS) in Fall 2021. There, I met many lasting collaborators and learned much of what I currently know about sampling. Overall, I found the program to be immensely intellectually stimulating and I am grateful.

I am indebted to collaborators and peers with whom I have had many fruitful conversations—unfortunately, too many to list here.

I am also grateful for the encouragement and suggestions from Sébastien Bubeck, Jonathan Niles-Weed, Philippe Rigollet, and Martin Wainwright, and for the careful reading of Michael Diao, Aram-Alexandre Pooladian, Sam Power, Yandi Shen, and Matthew S. Zhang.

Thank you to my family and friends who kept me sane throughout.

Part I

Diffusions in Continuous Time

The Langevin Diffusion in Continuous Time

In this chapter, we study the continuous-time **Langevin diffusion** with potential V , which is defined by the following stochastic differential equation (SDE):

$$dX_t = -\nabla V(X_t) dt + \sqrt{2} dB_t . \quad (1.0.1)$$

We begin with a brief primer on stochastic calculus in order to make sense of this equation. Then, we introduce two powerful frameworks for analyzing the Langevin diffusion: *Markov semigroup theory* and *the calculus of optimal transport*. These frameworks are two perspectives on the same diffusion, and the abstract calculus rules we develop within each framework streamline important computations.

A rigorous mathematical treatment of the theory in this chapter requires addressing substantial analytical technicalities, such as checking that the various partial differential equations (PDEs) are well-posed and that the calculations are carefully justified. We do not attempt to do so here and instead refer to the bibliography for detailed treatments. In particular, the “proofs” in this chapter are more akin to “proof sketches” which are meant to convey the main intuition.

1.1 A Primer on Stochastic Calculus

In the spirit of rapidly progressing to the study of sampling, the goal of this section is to familiarize the reader with the basic language of stochastic calculus, without delving into technical details. The mechanics of how to perform computations can be picked up through derivations in subsequent chapters. See [Steele \(2001\)](#); [Le Gall \(2016\)](#) for mathematical expositions. We treat further topics in stochastic calculus in Chapter 3.

1.1.1 The Itô Integral

In 1827, Robert Brown observed the curious jittery motion of particles of pollen in water. Subsequent study by mathematicians and physicists clarified the nature of the so-called “Brownian motion” as arising from collisions with surrounding molecules. The mathematical theory for this phenomenon

was initiated by the mathematician Louis Bachelier, who introduced a Brownian model in finance in 1900, as well as the physicist Albert Einstein in one of his breakthrough 1905 papers.

Since its inception, Brownian motion has been used to model the flow of heat, to price options in financial markets, to solve partial differential equations, to tease out the geometry of manifolds, and of course, to sample from probability distributions. It is perhaps not clear at first sight that such a process even exists,¹ but it would take us too far afield to give a construction here.

Definition 1.1.1. (Standard) **Brownian motion** is a stochastic process $(B_t)_{t \geq 0}$ in $\mathbb{R}^d$ satisfying the following properties:

- 1 $B_0 = 0$.
- 2 (Independence of increments) For all positive integers k and all $0 < t_1 < \dots < t_k$, the random variables $(B_{t_1}, B_{t_2} - B_{t_1}, \dots, B_{t_k} - B_{t_{k-1}})$ are mutually independent.
- 3 (Law of the increments) For all $0 \leq s < t < \infty$,

$$B_t - B_s \sim \text{normal}(0, (t - s) I_d).$$

- 4 (Continuity of the paths) Almost surely, $t \mapsto B_t$ is continuous.

For the rest of this subsection, we take the Brownian motion $(B_t)_{t \geq 0}$ to be one-dimensional. Our immediate goal is to develop the language of *stochastic calculus*, and in particular to define integrals involving Brownian motion: given a stochastic process $(\eta_t)_{t \geq 0}$ in $\mathbb{R}$, how do we make sense of an expression such as $\int_0^T \eta_t dB_t$? Once we have stochastic integration in hand, we can then formulate and solve *stochastic differential equations*.

The main technical difficulty in defining the stochastic integral is the irregularity of Brownian motion: by definition, $B_t \sim \text{normal}(0, t)$, so that $|B_t|$ is typically of size $\asymp \sqrt{t}$ for small $t > 0$. In particular, this means that Brownian motion is not differentiable at 0, or indeed, anywhere. Nevertheless, the stochastic integral can be meaningfully defined, and we summarize the main steps.

The standard measure-theoretic formalism works over a probability space $(\Omega, \mathcal{F}, \mathbb{P})$ which is *complete*, *filtered*, and *right-continuous*, meaning that there is an increasing family $(\mathcal{F}_t)_{t \geq 0}$ of σ -algebras with $\mathcal{F}_s \subseteq \mathcal{F}$ and $\bigcap_{t > s} \mathcal{F}_t = \mathcal{F}_s$ for all $s \geq 0$, and such that $\mathcal{F}_0$ contains all subsets of null sets. We assume that Brownian motion is adapted to the filtration: B_t is $\mathcal{F}_t$ -measurable for each $t \geq 0$. Although we make use of this formalism in this section to define the Itô integral, it is not important for the remainder of the book.

Defining the Itô integral at a single time T .

The procedure for defining the Itô integral is to work our way up, starting with simple processes and then extending the definition to richer ones. The simplest type of process is a piecewise constant one: $(\eta_t)_{t \geq 0}$ is of the form

$$\eta_t = \sum_{i=0}^{k-1} H_i \mathbb{1}_{\{t \in (t_i, t_{i+1}]\}} \quad (1.1.2)$$

¹Existence was established by Norbert Wiener, and hence the process is often called the *Wiener process*.

for some $0 \leq t_0 < t_1 < \dots < t_k$, where H_i is bounded and $\mathcal{F}_{t_i}$ -measurable. We call η an **elementary process**. In this case, perhaps the only reasonable definition of the stochastic integral is to take

$$\int_0^T \eta_t \, dB_t := \sum_{i=0}^{k-1} H_i (B_{t_{i+1} \wedge T} - B_{t_i \wedge T}). \quad (1.1.3)$$

This is indeed what we shall do, but for the moment we refrain from using the integral symbol and write this as $\mathcal{I}_{[0,T]}(\eta)$ to avoid confusion.

We record two key properties of the stochastic integral. The first is that $t \mapsto \mathcal{I}_{[0,t]}(\eta)$ is a continuous martingale, i.e., it is continuous and satisfies the following definition.

Definition 1.1.4. A process $(M_t)_{t \geq 0}$ is a **martingale** w.r.t. the filtration $(\mathcal{F}_t)_{t \geq 0}$ if for all $t \geq 0$, M_t is $\mathcal{F}_t$ -measurable and absolutely integrable, and

$$\mathbb{E}[M_t \mid \mathcal{F}_s] = M_s, \quad \text{for all } 0 \leq s < t.$$

Indeed, we deduce that $(\mathcal{I}_{[0,t]}(\eta))_{t \geq 0}$ is a martingale from two facts: (1) H_i is $\mathcal{F}_{t_i}$ -measurable for each i , and (2) $(B_t)_{t \geq 0}$ is a martingale.

The second key property is that we can compute the variance:

$$\mathbb{E}[\mathcal{I}_{[0,T]}(\eta)^2] = \mathbb{E}\left[\left|\sum_{i=0}^{k-1} H_i (B_{t_{i+1} \wedge T} - B_{t_i \wedge T})\right|^2\right] = \sum_{i=0}^{k-1} \mathbb{E}[|H_i (B_{t_{i+1} \wedge T} - B_{t_i \wedge T})|^2] \quad (1.1.5)$$

$$= \sum_{i=0}^{k-1} \mathbb{E}[H_i^2] ((t_{i+1} \wedge T) - (t_i \wedge T)) = \mathbb{E} \int_0^T \eta_t^2 \, dt. \quad (1.1.6)$$

Here, we used the basic properties listed in the definition of Brownian motion, such as independence of increments. This equation shows that if $\mathbb{P}_T := \mathbb{P} \otimes \mathfrak{m}|_{[0,T]}$, where $\mathfrak{m}|_{[0,T]}$ is the Lebesgue measure on $[0, T]$, then the mapping $\eta \mapsto \mathcal{I}_{[0,T]}(\eta)$ is an isometry, from the subspace of $L^2(\mathbb{P}_T)$ consisting of elementary processes, to $L^2(\mathbb{P})$. Through this isometry, we can extend the definition of the stochastic integral to the closure of the elementary processes.

The closure turns out to consist of processes $(\eta_t)_{t \geq 0}$ that are *progressive*² and satisfies the integrability condition

$$\|\eta\|_{L^2(\mathbb{P}_T)}^2 = \mathbb{E} \int_0^T \eta_t^2 \, dt < \infty. \quad (1.1.7)$$

This means that $(\eta_t)_{t \geq 0}$ can be approximated by elementary processes $\{(\eta_t^{(k)})_{t \geq 0} : k \in \mathbb{N}\}$ of the form (1.1.2) in the $L^2(\mathbb{P}_T)$ norm. For each k , the stochastic integral $\mathcal{I}_{[0,T]}(\eta^{(k)})$ is defined via (1.1.3), and $\lim_{k \rightarrow \infty} \mathcal{I}_{[0,T]}(\eta^{(k)})$ exists in $L^2(\mathbb{P})$ thanks to the isometry. We can then take the limit to be the definition of the stochastic integral $\mathcal{I}_{[0,T]}(\eta)$.

Defining the Itô integral as a stochastic process.

Although the procedure above successfully defines $\mathcal{I}_{[0,t]}(\eta)$ for a *fixed* time $t > 0$, there is no guarantee of coherence between different times t . The trouble arises because $\mathcal{I}_{[0,t]}(\eta)$ is defined as a limit, but this limit is only well-specified up to an event of measure zero, and these measure zero events for

²The process $(\eta_t)_{t \geq 0}$ is progressive if for all $T \geq 0$, the mapping $(\omega, t) \mapsto \eta_t(\omega)$ is measurable w.r.t. $\mathcal{F}_T \otimes \mathcal{B}_{[0,T]}$, where $\mathcal{B}_{[0,T]}$ is the Borel σ -algebra on $[0, T]$.

different times t might conceivably accumulate into something more. This is undesirable because the true power of stochastic calculus comes from viewing the stochastic integral as a time-indexed stochastic process in its own right.

The key insight, due to Kiyosi Itô, is to go back to the approximating sequence $\{(\eta_t^{(k)})_{t \geq 0} : k \in \mathbb{N}\}$. For each k , the Itô integral is defined as an entire process $t \mapsto \mathcal{I}_{[0,t]}(\eta^{(k)})$ via (1.1.3), and moreover this process is a continuous martingale. One can then apply powerful results on martingale convergence, leading to the following theorem.

Theorem 1.1.8 (Definition of the Itô integral; Itô isometry). *Suppose that $(\eta_t)_{t \in [0,T]}$ is progressive and satisfies $\mathbb{E} \int_0^T \eta_t^2 dt < \infty$. Then, there exists a continuous martingale, denoted $(\int_0^t \eta_s dB_s)_{t \in [0,T]}$, which is adapted to $(\mathcal{F}_t)_{t \in [0,T]}$ and satisfies*

$$\mathbb{E} \left[\left| \int_0^t \eta_s dB_s \right|^2 \right] = \mathbb{E} \int_0^t \eta_s^2 ds, \quad \text{for all } t \in [0, T]. \quad (1.1.9)$$

*The formula (1.1.9) is called the **Itô isometry**.*

Also, for each $t \in [0, T]$, it holds that $\int_0^t \eta_s dB_s = \mathcal{I}_{[0,t]}(\eta)$ a.s.

Extending the definition via localization.

There is one final step which is typically taken, namely to expand the class of allowable integrands to progressive processes η with

$$\int_0^T \eta_s^2 ds < \infty \quad \text{almost surely.} \quad (1.1.10)$$

Note that this condition is weaker than the condition $\mathbb{E} \int_0^T \eta_s^2 ds < \infty$. Such an extension is evidently mathematically interesting, as we would like our definitions to be as broad as possible. However, equally important is that it introduces the device of **localization**. On the whole, localization actually serves to *reduce* the number of technicalities in the subject: once introduced, it allows us to always work with a stopping time up to which the process is as nice as one desires (e.g., bounded). Localization is also needed to state the most useful version of Itô's formula in Theorem 1.1.19. However, this is not the focus of the book, and in what follows we usually brush over such localization arguments. For now, we simply introduce the definitions in order to convey that the idea is not too complicated.

Definition 1.1.11. A **stopping time** τ is a random variable such that for each $t \geq 0$, the event $\{\tau \leq t\}$ is $\mathcal{F}_t$ -measurable.

Definition 1.1.12. An increasing sequence of stopping times $(\tau_n)_{n \in \mathbb{N}}$ is called a **localizing sequence** for η on $[0, T]$ if:

- 1 for all $n \in \mathbb{N}$, $(\eta_t \mathbb{1}_{\{t \leq \tau_n\}})_{t \geq 0}$ has finite $\|\cdot\|_{L^2(\mathbb{P}_T)}$ norm, and
- 2 $\tau_n \rightarrow T$ almost surely.

The good news is that localizing sequences are easy to find, and the following proposition barely needs a proof (and so we omit it).

Proposition 1.1.13. *If η is a progressive process satisfying the condition (1.1.10), then the sequence $(\tau_n)_{n \in \mathbb{N}}$ defined by*

$$\tau_n := \inf \left\{ t \geq 0 \mid \int_0^t \eta_s^2 ds \geq n \right\} \wedge T$$

is a localizing sequence for η on $[0, T]$.

The idea now is straightforward, even if the execution is not: for each progressive process η satisfying (1.1.10), let $(\tau_n)_{n \in \mathbb{N}}$ be a localizing sequence for η on $[0, T]$. For each $n \in \mathbb{N}$, by the definition of a localizing sequence, we can apply our existing definition of the Itô integral to the stopped process $(\eta_t \mathbb{1}_{t \leq \tau_n})_{t \geq 0}$, which gives us a continuous martingale

$$t \mapsto \int_0^t \eta_s \mathbb{1}_{\{s \leq \tau_n\}} dB_s. \quad (1.1.14)$$

Then, we can define the Itô integral of η to be a limit of the processes (1.1.14). The details are technical, and omitted.

We can also define an analogue of martingales using localizing sequences; these are almost martingales, but lack the required integrability.

Definition 1.1.15. A process $(M_t)_{t \geq 0}$ is a **local martingale** if it is adapted to the filtration $(\mathcal{F}_t)_{t \geq 0}$ and there is an increasing sequence $(\tau_n)_{n \in \mathbb{N}}$ of stopping times such that $\tau_n \rightarrow \infty$ and for each n , the process $t \mapsto M_{t \wedge \tau_n} - M_0$ is a martingale w.r.t. $(\mathcal{F}_t)_{t \geq 0}$.

Proposition 1.1.16. *If η is a progressive process satisfying (1.1.10), then the Itô integral $t \mapsto \int_0^t \eta_s dB_s$ is a continuous local martingale.*

As an example, consider $\eta_t = \exp(B_t^2)$ for $t \geq 0$. For all $T > 1/4$, this does not satisfy the condition (1.1.7), but we can still define the Itô integral $\int_0^T \eta_t dB_t$ using localizing sequences. In general, the key advantage of localization is that we do not have to check conditions such as (1.1.7) *a priori* in order to make sense of our computations.

Looking forward.

The construction of the Itô integral may seem rather abstract, and we are sorely lacking in examples. At this juncture, it is common to work out simple exercises such as computing $\int_0^t B_s dB_s$, and while this is pedagogically natural, it is also liable to mislead the reader into thinking that the main use of Itô integration is to solve synthetic problems with no apparent purpose. Although counterintuitive, our answer to the heavy amount of abstraction is *more abstraction*. In the next subsection, we develop the single most important computation rule in stochastic calculus (along with the Itô isometry (1.1.9)), called Itô's formula, after which we hardly need to return to the definition of a stochastic integral ever again. And even Itô's formula will be abstracted out into the language of Markov semigroups in Section 1.2. The upshot is that we introduced the Itô integral because it is the foundation of our field, but the details of what we have developed thus far are not needed for the remainder of the book.

1.1.2 Itô's Formula

With the Itô integral in hand, we consider the following class of processes.

Definition 1.1.17. A stochastic process $(X_t)_{t \geq 0}$ is an **Itô process** if it is of the form

$$X_t = X_0 + \int_0^t b_s \, ds + \int_0^t \sigma_s \, dB_s, \quad \text{for } t \geq 0,$$

where $(b_t)_{t \geq 0}$ takes values in $\mathbb{R}^d$, $(\sigma_t)_{t \geq 0}$ takes values in $\mathbb{R}^{d \times N}$, and $(B_t)_{t \geq 0}$ is a standard Brownian motion in $\mathbb{R}^N$.

Implicit in the above definition is that the process should be well-defined: the coefficients $(b_t)_{t \geq 0}$ and $(\sigma_t)_{t \geq 0}$ should be progressive processes for which the integrals exist. In other words, $\int_0^t \|b_s\| \, ds < \infty$ and $\int_0^t \|\sigma_s\|_{\text{HS}}^2 \, ds < \infty$ almost surely, where we write $\|M\|_{\text{HS}}$ for the Hilbert–Schmidt (or Frobenius) norm of M , given by

$$\|M\|_{\text{HS}}^2 := \text{tr}(MM^\top) = \text{tr}(M^\top M).$$

Also, the random variable X_0 should be $\mathcal{F}_0$ -measurable, and then the process $(X_t)_{t \geq 0}$ is also progressive.

We refer to $(b_t)_{t \geq 0}$ as the **drift coefficient** and $(\sigma_t)_{t \geq 0}$ as the **diffusion coefficient**. When the drift coefficient is zero, then $(X_t)_{t \geq 0}$ is simply an Itô integral, and thus a continuous local martingale (Proposition 1.1.16). Otherwise, for a non-zero drift coefficient, the process $(X_t)_{t \geq 0}$ is no longer necessarily a local martingale. As a shorthand, we often write the Itô process in differential form:

$$dX_t = b_t \, dt + \sigma_t \, dB_t. \quad (1.1.18)$$

Our goal is to understand how the Itô process transforms when we compose it with a smooth function $f : \mathbb{R}^d \rightarrow \mathbb{R}$. This leads to *Itô's formula*, which is the main engine of stochastic calculus.

Although the notation (1.1.18) is informal, it conveys the main intuition. For $h > 0$ small, we can approximate $X_{t+h} \approx X_t + h b_t + \sqrt{h} \sigma_t \xi$, where $\xi \sim \text{normal}(0, I_N)$. Note that the $\sqrt{h}$ scaling comes from the fact that the Brownian increment $B_{t+h} - B_t$ has the $\text{normal}(0, hI_N)$ distribution. Now suppose that $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is twice continuously differentiable. Normally, to compute $f(X_{t+h}) - f(X_t)$ up to order $o(h)$, a *first-order* Taylor expansion of f suffices, but in stochastic calculus this would miss important terms arising from the Brownian motion: indeed, second-order terms in $B_{t+h} - B_t$ are of order h and hence not negligible.

Therefore, we carry out the Taylor expansion to an extra term:

$$\begin{aligned} f(X_{t+h}) - f(X_t) &\approx \langle \nabla f(X_t), h b_t + \sqrt{h} \sigma_t \xi \rangle + \frac{1}{2} \langle h b_t + \sqrt{h} \sigma_t \xi, \nabla^2 f(X_t) (h b_t + \sqrt{h} \sigma_t \xi) \rangle \\ &= h \left\{ \langle \nabla f(X_t), b_t \rangle + \frac{1}{2} \langle \sigma_t \xi, \nabla^2 f(X_t) \sigma_t \xi \rangle \right\} + \sqrt{h} \langle \sigma_t^\top \nabla f(X_t), \xi \rangle + o(h). \end{aligned}$$

This expression suggests that $(f(X_t))_{t \geq 0}$ is also an Itô process. The third term, which is of order $\sqrt{h}$, turns into an Itô integral once integrated. Perhaps the most interesting term is the second term, $\frac{h}{2} \langle \sigma_t^\top \nabla^2 f(X_t) \sigma_t, \xi \xi^\top \rangle$, which is a genuinely new feature of stochastic calculus. If we sum up many of these increments, we end up with an expression like $\frac{1}{2} \sum_{k=0}^K \langle \sigma_{t+kh}^\top \nabla^2 f(X_{t+kh}) \sigma_{t+kh}, h \xi_k \xi_k^\top \rangle$. If we replace each $h \xi_k \xi_k^\top$ by its expectation $h I_N$ (which must be carefully justified; see the calculation of quadratic variation in Section 3.1), then this resembles a Riemann sum, which converges to the integral of $\frac{1}{2} \langle \nabla^2 f(X_t), \sigma_t \sigma_t^\top \rangle$. This is formalized in the following theorem.

Theorem 1.1.19 (Itô's formula). *Let $(X_t)_{t \geq 0}$ be an Itô process, $dX_t = b_t dt + \sigma_t dB_t$, and let $f \in C^2(\mathbb{R}^d)$. Then, $(f(X_t))_{t \geq 0}$ is also an Itô process which satisfies, for $t \geq 0$:*

$$f(X_t) - f(X_0) = \int_0^t \left\{ \langle \nabla f(X_s), b_s \rangle + \frac{1}{2} \langle \nabla^2 f(X_s), \sigma_s \sigma_s^\top \rangle \right\} ds + \int_0^t \langle \sigma_s^\top \nabla f(X_s), dB_s \rangle.$$

We omit the proof, since the bulk of the intuition is carried in the informal Taylor series argument described above. Observe that since Itô integrals are (under appropriate integrability conditions) continuous martingales, the expectation of the last term in Itô's formula is typically zero. Therefore,

$$\mathbb{E} f(X_t) - \mathbb{E} f(X_0) = \int_0^t \mathbb{E} \left[\langle \nabla f(X_s), b_s \rangle + \frac{1}{2} \langle \nabla^2 f(X_s), \sigma_s \sigma_s^\top \rangle \right] ds, \quad (1.1.20)$$

or in differential form,

$$\partial_t \mathbb{E} f(X_t) = \mathbb{E} \left[\langle \nabla f(X_t), b_t \rangle + \frac{1}{2} \langle \nabla^2 f(X_t), \sigma_t \sigma_t^\top \rangle \right].$$

Itô's formula can also be extended to time-dependent functions via

$$\begin{aligned} f(t, X_t) - f(0, X_0) &= \int_0^t \left\{ \partial_s f(s, X_s) + \langle \nabla f(s, X_s), b_s \rangle + \frac{1}{2} \langle \nabla^2 f(s, X_s), \sigma_s \sigma_s^\top \rangle \right\} ds \\ &\quad + \int_0^t \langle \sigma_s^\top \nabla f(s, X_s), dB_s \rangle. \end{aligned}$$

We revisit and streamline Itô's formula in Section 3.1.

1.1.3 Existence and Uniqueness of SDEs

Let $b : \mathbb{R}_+ \times \mathbb{R}^d \rightarrow \mathbb{R}^d$ and $\sigma : \mathbb{R}_+ \times \mathbb{R}^d \rightarrow \mathbb{R}^{d \times N}$. Consider the stochastic differential equation (SDE)

$$dX_t = b(t, X_t) dt + \sigma(t, X_t) dB_t. \quad (1.1.21)$$

Suppose we are given a complete filtered probability space $(\Omega, \mathcal{F}, (\mathcal{F}_t)_{t \geq 0}, \mathbb{P})$ which supports a standard N -dimensional adapted Brownian motion B . A **(strong) solution** to the SDE is an adapted $\mathbb{R}^d$ -valued process X such that

$$X_t = X_0 + \int_0^t b(s, X_s) ds + \int_0^t \sigma(s, X_s) dB_s.$$

In this section, we state the basic existence and uniqueness result for this SDE. The result states that if the coefficients b, σ are Lipschitz in space uniformly in time, then the SDE admits a unique solution. Uniqueness means that if there are two solutions $X, \tilde{X}$ to SDE on the same probability space, driven by the same Brownian motion, then $X = \tilde{X}$.

Theorem 1.1.22 (Existence and uniqueness of SDE solutions). *Assume that b and σ are continuous, and there exists $C > 0$ such that for all $t \geq 0$ and $x, y \in \mathbb{R}^d$,*

$$\|b(t, x) - b(t, y)\| \vee \|\sigma(t, x) - \sigma(t, y)\|_{\text{HS}} \leq C \|x - y\|.$$

Then, for every $T > 0$, $x \in \mathbb{R}^d$, and complete filtered probability space $(\Omega, \mathcal{F}, (\mathcal{F}_t)_{t \geq 0}, \mathbb{P})$ supporting the adapted Brownian motion B , there exists a pathwise unique solution $(X_t)_{t \in [0, T]}$ for the SDE (1.1.21) with $X_0 = x$. Moreover, the solution $(X_t)_{t \in [0, T]}$ is a Markov process.

A theorem similar in spirit to [Theorem 1.1.22](#) can be established under the assumption that b and σ are only *locally* Lipschitz, but in this case the solution to the SDE is not guaranteed to exist for all time. The issue is that when the coefficients grow faster than linearly, there can be a finite (random) time $\mathfrak{e}$, called the *explosion time*, such that $\|X_t\| \rightarrow \infty$ as $t \rightarrow \mathfrak{e}$. This phenomenon is already present for ODEs (see [Exercise 1.2](#)). For our purposes, the assumption of Lipschitz coefficients suffices.

1.2 Markov Semigroup Theory

More thorough treatments of Markov semigroup theory can be found in [Bakry et al. \(2014\)](#); [van Handel \(2016\)](#). We will revisit and expand upon many of the topics introduced here in Chapter 2.

1.2.1 Basic Definitions and Kolmogorov's Equations

The core idea of Markov semigroup theory is to encode the behavior of a Markov process $(X_t)_{t \geq 0}$ via operators which act on functions. We will then develop calculus rules for working with these operators, and we can study the operators via functional analysis. This is analogous to how the linear algebraic study of the transition matrix of a discrete-time Markov chain reveals properties (e.g., ergodicity, convergence) of the chain.

Definition 1.2.1. For a time-homogeneous Markov process $(X_t)_{t \geq 0}$, its associated **Markov semigroup** $(P_t)_{t \geq 0}$ is the family of operators acting on functions via

$$P_t f(x) := \mathbb{E}[f(X_t) \mid X_0 = x].$$

The Markov property and iterated conditioning yield the following lemma (exercise).

Lemma 1.2.2 (Semigroup property). *The Markov semigroup $(P_t)_{t \geq 0}$ satisfies $P_0 = \text{id}$ and $P_s P_t = P_t P_s = P_{s+t}$ for all $s, t \geq 0$.*

The starting point of developing a calculus is to differentiate the semigroup $t \mapsto P_t$, which is accomplished via the following definition.

Definition 1.2.3. The **infinitesimal generator** $\mathcal{L}$ associated with a Markov semigroup $(P_t)_{t \geq 0}$ is the operator defined by

$$\mathcal{L}f := \lim_{t \searrow 0} \frac{P_t f - f}{t},$$

for all functions f for which the above limit exists.

Here we pause to warn the reader of some technical issues. Mathematical difficulties of Markov semigroup theory arise in trying to answer the following questions: on what space of functions is the generator defined, and in what sense is the above limit taken? As we shall see, a natural space of functions to consider is $L^2(\pi)$, with π denoting the stationary distribution of the diffusion. However, the generator is usually a differential operator, and not all functions in $L^2(\pi)$ have enough regularity to lie in the domain of the generator. The theory of *unbounded* linear operators on a Hilbert space was developed to handle this situation, but it is rife with subtle distinctions such as the difference between symmetric and self-adjoint operators. We will brush over these issues, and all generator identities below are to be understood over the relevant domains on which the calculations are justified; see [Bakry et al. \(2014, Chapter 3\)](#) for details.

Example 1.2.4. Let us compute the generator of the Langevin diffusion (1.0.1). In fact, the following computation is simply a consequence of Itô's formula ([Theorem 1.1.19](#)), but it does not hurt to derive this result from scratch. We approximate

$$X_t = X_0 - \int_0^t \nabla V(X_s) ds + \sqrt{2} B_t = X_0 - t \nabla V(X_0) + \sqrt{2} B_t + o(t).$$

Assuming that $f \in C^2(\mathbb{R}^d)$ with bounded derivatives, we perform a Taylor expansion of f to second order.

$$\mathbb{E} f(X_t) = \mathbb{E}[f(X_0) + \langle \nabla f(X_0), -t \nabla V(X_0) + \sqrt{2} B_t \rangle + \langle \nabla^2 f(X_0) B_t, B_t \rangle] + o(t).$$

Since B_t is mean zero and independent of X_0 , with $\mathbb{E}[B_t B_t^\top] = t I_d$,

$$\begin{aligned} \mathbb{E}[f(X_t) \mid X_0 = x] &= f(x) - t \langle \nabla f(x), \nabla V(x) \rangle + t \operatorname{tr} \nabla^2 f(x) + o(t) \\ &= f(x) - t \langle \nabla f(x), \nabla V(x) \rangle + t \Delta f(x) + o(t). \end{aligned}$$

Hence,

$$\mathcal{L}f(x) = \lim_{t \searrow 0} \frac{\mathbb{E}[f(X_t) \mid X_0 = x] - f(x)}{t} = \Delta f(x) - \langle \nabla V(x), \nabla f(x) \rangle.$$

The Markov semigroup and dynamics.

As promised, the Markov semigroup captures the information that was contained in the original Markov process. One way to demonstrate this is to prove theorems which show that the Markov process can be completely recovered from its Markov semigroup. Another approach, which we now take up, is to show that the dynamics of the Markov process are captured via calculation rules involving the Markov semigroup.

Proposition 1.2.5 (Kolmogorov's backward equation). *For all $t \geq 0$, it holds that*

$$\partial_t P_t f = \mathcal{L} P_t f = P_t \mathcal{L} f .$$

In particular, $\mathcal{L}$ commutes with the semigroup $(P_t)_{t \geq 0}$.

Proof Observe that

$$\lim_{h \searrow 0} \frac{P_{t+h} f - P_t f}{h} = \lim_{h \searrow 0} \frac{P_h - \text{id}}{h} P_t f = \mathcal{L} P_t f .$$

Repeating the computation but factoring out P_t on the left yields the second equality. $\square$

There is a dual to this equation: let π_0 denote the density of X_0 . Formally, we can write $\mathbb{E} f(X_t) = \int P_t f d\pi_0 = \int f dP_t^* \pi_0$, where P_t^* is the adjoint of P_t . This says that the law π_t of X_t is formally given by $P_t^* \pi_0$. Moreover, by Kolmogorov's backward equation,

$$\partial_t \int f dP_t^* \pi_0 = \partial_t \int P_t f d\pi_0 = \int P_t \mathcal{L} f d\pi_0 = \int \mathcal{L} f dP_t^* \pi_0 = \int f d\mathcal{L}^* P_t^* \pi_0 .$$

Since this has to hold for all functions f , we conclude the following.

Proposition 1.2.6 (Kolmogorov's forward equation). *For all $t \geq 0$,*

$$\partial_t P_t^* \pi_0 = \mathcal{L}^* P_t^* \pi_0 = P_t^* \mathcal{L}^* \pi_0 .$$

Here is another illuminating way to express these equations. For $u_t := P_t f$, and $\pi_t = P_t^* \pi_0$,

$$\partial_t u_t = \mathcal{L} u_t , \quad (\text{Kolmogorov's backward equation})$$

$$\partial_t \pi_t = \mathcal{L}^* \pi_t . \quad (\text{Kolmogorov's forward equation})$$

The terms “backward” and “forward” can evoke confusion, so we will not use them. Instead, we will henceforth refer to the evolution equation for the density (Kolmogorov's forward equation) as the **Fokker–Planck equation**.

Consequently, we obtain characterizations of stationarity. Recall that π is stationary for the Markov process if, when $X_0 \sim \pi$, then $X_t \sim \pi$ for all $t \geq 0$.

Proposition 1.2.7 (Stationarity). *The following are equivalent.*

- 1 π is a stationary distribution for the Markov process.
- 2 $\mathcal{L}^* \pi = 0$.
- 3 $\mathbb{E}_\pi \mathcal{L} f = 0$ for all functions f .

Proof The equivalence between the first two statements is the Fokker–Planck equation. The third statement is the dual of the second statement. $\square$

Example 1.2.8. Consider again the Langevin diffusion. For functions $f, g : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\int \mathcal{L} f g = \int \{\Delta f - \langle \nabla V, \nabla f \rangle\} g = \int f \{\Delta g + \operatorname{div}(g \nabla V)\}$$

where the second equality is integration by parts. This shows that

$$\mathcal{L}^* g = \Delta g + \operatorname{div}(g \nabla V).$$

From here, we can solve for the stationary distribution. Write

$$0 = \mathcal{L}^* \pi = \Delta \pi + \operatorname{div}(\pi \nabla V) = \operatorname{div}(\pi (\nabla \log \pi + \nabla V)).$$

This is solved by setting $\log \pi = -V + \text{constant}$, i.e., $\pi \propto \exp(-V)$, provided $\int \exp(-V) < \infty$.

Corollary 1.2.9. Under regularity suitable assumptions, the unique stationary distribution of the Langevin diffusion (1.0.1) with potential V is $\pi \propto \exp(-V)$.

1.2.2 Reversibility and the Spectrum

Consider a Markov semigroup $(P_t)_{t \geq 0}$ with generator $\mathcal{L}$ and stationary distribution π . Then, the natural space of functions to study is the Hilbert space $L^2(\pi)$. The analysis of the Markov process is particularly simple if the following condition holds.

Definition 1.2.10. The Markov semigroup $(P_t)_{t \geq 0}$ is **reversible** w.r.t. π if for all $f, g \in L^2(\pi)$ and all $t \geq 0$,

$$\int P_t f g \, d\pi = \int f P_t g \, d\pi.$$

Equivalently, for all f and g for which $\mathcal{L}f$ and $\mathcal{L}g$ are defined,

$$\int \mathcal{L} f g \, d\pi = \int f \mathcal{L} g \, d\pi.$$

If $X_0 \sim \pi$ and we take $f = \mathbb{1}_A$ and $g = \mathbb{1}_B$ for events A and B , then it implies

$$\mathbb{P}\{X_t \in A, X_0 \in B\} = \mathbb{P}\{X_0 \in A, X_t \in B\},$$

i.e., (X_0, X_t) has the same distribution as (X_t, X_0) . This is the sense in which the associated Markov process is time-reversible.

The definition says that P_t and $\mathcal{L}$ are symmetric operators on $L^2(\pi)$, and thus we expect that P_t and $\mathcal{L}$ have real spectra. Also, since $\partial_t P_t = \mathcal{L} P_t$, we can formally write $P_t = \exp(t\mathcal{L})$, and so we expect P_t to be a positive operator, meaning that $\int f P_t f \, d\pi \geq 0$ for all $f \in L^2(\pi)$, which can be checked from reversibility (write $\int f P_t f \, d\pi = \int (P_{t/2} f)^2 \, d\pi$). Moreover, from the definition of P_t and Jensen's inequality,

$$\{P_t f(x)\}^2 = \mathbb{E}[f(X_t) \mid X_0 = x]^2 \leq \mathbb{E}[f(X_t)^2 \mid X_0 = x] = P_t(f^2)(x) \quad (1.2.11)$$

and integrating this yields $\int (P_t f)^2 \, d\pi \leq \int P_t(f^2) \, d\pi = \int f^2 \, d\pi$, where the equality follows from

stationarity of π . This shows that P_t is a (non-strict) contraction on $L^2(\pi)$ (in fact, on any $L^p(\pi)$, $p \in [1, \infty]$). Combining this with $P_t = \exp(t\mathcal{L})$ leads us to predict that $\mathcal{L}$ is a *negative* operator. Below, we will give a direct proof of this fact; unsurprisingly, the proof of negativity of $\mathcal{L}$ still relies on the crucial fact (1.2.11).

Definition 1.2.12. The **carré du champ** is the bilinear operator Γ defined via

$$\Gamma(f, g) := \frac{1}{2} \{ \mathcal{L}(fg) - f \mathcal{L}g - g \mathcal{L}f \}.$$

The **Dirichlet energy** is the functional $\mathcal{E}(f, g) := \int \Gamma(f, g) d\pi$.

Lemma 1.2.13 (Non-negativity of carré du champ). *For any function f , $\Gamma(f, f) \geq 0$.*

Proof Recall from (1.2.11) that $P_t(f^2) \geq (P_t f)^2$ for all $t > 0$. In terms of $\mathcal{L}$,

$$f^2 + t \mathcal{L}(f^2) + o(t) \geq [f + t \mathcal{L}f + o(t)]^2 = f^2 + 2t f \mathcal{L}f + o(t)$$

and sending $t \searrow 0$ yields $\mathcal{L}(f^2) \geq 2f \mathcal{L}f$. (This proof does not require reversibility.) $\square$

Theorem 1.2.14 (Fundamental integration by parts identity). *Suppose that the generator $\mathcal{L}$ and carré du champ Γ are associated with a Markov semigroup which is reversible w.r.t. π . Then, for any functions f and g ,*

$$\int f (-\mathcal{L})g d\pi = \int (-\mathcal{L})f g d\pi = \int \Gamma(f, g) d\pi = \mathcal{E}(f, g).$$

Since the identity implies that $\mathcal{L}$ is symmetric, the integration by parts identity is in fact *equivalent* to reversibility. It also implies $\int f (-\mathcal{L})f d\pi \geq 0$ for all f .

Corollary 1.2.15. *For a reversible Markov semigroup, $-\mathcal{L} \geq 0$.*

Proof of Theorem 1.2.14 Since $\int \mathcal{L}h d\pi = 0$ for all functions h (due to stationarity of π), the definition of Γ yields

$$\int \Gamma(f, g) d\pi = \frac{1}{2} \int f (-\mathcal{L})g d\pi + \frac{1}{2} \int g (-\mathcal{L})f d\pi.$$

The two terms are equal due to reversibility. $\square$

It is usually convenient for our operators to be positive, so from now on we will instead refer to the negative generator $-\mathcal{L}$.

When we introduced Kolmogorov's equations, we ended up with two PDEs, one involving the generator $\mathcal{L}$ and one involving its $L^2(\mathfrak{m})$ adjoint $\mathcal{L}^*$, where $\mathfrak{m}$ is the Lebesgue measure on $\mathbb{R}^d$. The issue is that we used the “wrong” inner product; $\mathcal{L}$ is not symmetric in $L^2(\mathfrak{m})$. If we now switch to $L^2(\pi)$, then instead of considering the density π_t with respect to $\mathfrak{m}$ we should consider the density

$\rho_t := \pi_t/\pi$ with respect to π . Then, the Fokker–Planck equation becomes

$$\partial_t \rho_t = \mathcal{L} \rho_t. \quad (1.2.16)$$

For the rest of the section, the Markov semigroup is assumed reversible.

Example 1.2.17. Returning to the fundamental example of the Langevin diffusion, a computation shows that

$$\Gamma(f, f) = \frac{1}{2} \{ \Delta(f^2) - \langle \nabla V, \nabla(f^2) \rangle - 2f (\Delta f - \langle \nabla V, \nabla f \rangle) \} = \|\nabla f\|^2.$$

Incidentally, *carré du champ* means “square of the field” in French, and it is this expression which gives it its name. More generally, $\Gamma(f, g) = \langle \nabla f, \nabla g \rangle$.

The identity in Theorem 1.2.14 reads

$$\int f (-\Delta g + \langle \nabla V, \nabla g \rangle) d\pi = \int g (-\Delta f + \langle \nabla V, \nabla f \rangle) d\pi = \int \langle \nabla f, \nabla g \rangle d\pi$$

which can be checked using $\pi \propto \exp(-V)$ via integration by parts (naturally!), showing that the Langevin diffusion is indeed reversible w.r.t. π .

We can reformulate Theorem 1.2.14 in this case in an illuminating way as

$$-\mathcal{L} = \nabla^{*,\pi} \nabla \quad (1.2.18)$$

where $(\cdot)^{*,\pi}$ denotes the adjoint in $L^2(\pi)$. This expression brings out the symmetry of $\mathcal{L}$.

Gradient flow of the Dirichlet energy.

It turns out that a reversible Markov process follows the *steepest descent* of the Dirichlet energy with respect to $L^2(\pi)$. To justify this, for a curve $t \mapsto u_t$ in $L^2(\pi)$, write $\dot{u}_t := \partial_t u_t$ for the time derivative. The $L^2(\pi)$ gradient of the functional $f \mapsto \mathcal{E}(f) := \mathcal{E}(f, f)$ at f is defined to be the element $\nabla_{L^2(\pi)} \mathcal{E}(f) \in L^2(\pi)$ such that for all curves $t \mapsto u_t$ with $u_0 = f$, it holds that

$$\partial_t \Big|_{t=0} \mathcal{E}(u_t, u_t) = \int \dot{u}_0 \nabla_{L^2(\pi)} \mathcal{E}(f) d\pi.$$

From the integration by parts identity,

$$\partial_t \Big|_{t=0} \mathcal{E}(u_t, u_t) = \partial_t \Big|_{t=0} \int u_t (-\mathcal{L}) u_t d\pi = 2 \int \dot{u}_0 (-\mathcal{L}) f d\pi.$$

Therefore, $\nabla_{L^2(\pi)} \mathcal{E}(f) = -2\mathcal{L}f$.

The steepest descent of $\mathcal{E}$ is the curve $t \mapsto u_t$ such that $\dot{u}_t = -\nabla_{L^2(\pi)} \mathcal{E}(u_t) = 2\mathcal{L}u_t$. This is, up to a rescaling of time, precisely the equation satisfied by $t \mapsto P_t f$.

Spectral gap and convergence.

Consider a reversible Markov semigroup $(P_t)_{t \geq 0}$ and recall that $P_t f(x) = \mathbb{E}[f(X_t) \mid X_0 = x]$. We are interested in the long-term behavior of $P_t f$. If the process mixes, then by definition it forgets its initial condition, so that $P_t f$ converges to a constant; moreover, this constant should be the average value $\int f d\pi$ at stationarity. How do we establish a rate of convergence for $P_t f \rightarrow \int f d\pi$?

We may assume that $\int f d\pi = 0$, so we wish to prove $P_t f \rightarrow 0$. Also recall that, formally,

$P_t = \exp(t\mathcal{L})$ with $\mathcal{L} \leq 0$. If we have a *spectral gap*

$$-\mathcal{L} \geq \lambda_{\min} > 0,$$

then we would expect that

$$\|P_t f\|_{L^2(\pi)}^2 \leq \exp(-2\lambda_{\min} t) \|f\|_{L^2(\pi)}^2.$$

This is indeed the case. However, observe that since $P_t 1 = 1$ for all $t \geq 0$ and hence $\mathcal{L}1 = 0$, the spectral gap condition is only supposed to hold on the subspace $1^\perp \subseteq L^2(\pi)$ which is orthogonal to the space of constant functions.

Definition 1.2.19. The Markov process is said to satisfy a **Poincaré inequality** (PI) with constant C_{PI} if for all functions $f \in L^2(\pi)$,

$$\int f (-\mathcal{L}) f \, d\pi = \mathcal{E}(f, f) \geq \frac{1}{C_{\text{PI}}} \|\text{proj}_{1^\perp} f\|_{L^2(\pi)}^2 = \frac{1}{C_{\text{PI}}} \text{var}_\pi f.$$

The Poincaré constant C_{PI} corresponds to the inverse of the spectral gap. Under the assumption of a Poincaré inequality, we can establish convergence in $L^2(\pi)$ via Markov semigroup calculus:

$$\begin{aligned} \partial_t \|P_t f\|_{L^2(\pi)}^2 &= \partial_t \int (P_t f)^2 \, d\pi = 2 \int P_t f \mathcal{L} P_t f \, d\pi = -2\mathcal{E}(P_t f, P_t f) \\ &\leq -\frac{2}{C_{\text{PI}}} \|P_t f\|_{L^2(\pi)}^2. \end{aligned}$$

Such differential inequalities are encountered repeatedly throughout the book and are handled with Grönwall's inequality.

Lemma 1.2.20 (Grönwall's lemma, differential form). *Let $T > 0$ and let $g : [0, T] \rightarrow \mathbb{R}$ be differentiable. Suppose that there is a constant $c \in \mathbb{R}$ such that*

$$g'(t) \leq c g(t), \quad \forall t \in [0, T].$$

Then,

$$g(t) \leq g(0) \exp(ct), \quad \forall t \in [0, T].$$

Proof Differentiate the function $t \mapsto \exp(-ct) g(t)$ and apply the assumption.

$$\partial_t [\exp(-ct) g(t)] = \exp(-ct) [-c g(t) + g'(t)] \leq 0. \quad \square$$

Hence, the differential inequality above proves the implication from the first to the second statement below. The converse implication is left as [Exercise 1.6](#).

Theorem 1.2.21 (PI and variance decay). *The following are equivalent.*

- 1 *The Markov process satisfies a Poincaré inequality with constant C_{PI} .*
- 2 *For all $f \in L^2(\pi)$ with $\int f \, d\pi = 0$ and all $t \geq 0$,*

$$\|P_t f\|_{L^2(\pi)}^2 \leq \exp\left(-\frac{2t}{C_{\text{PI}}}\right) \|f\|_{L^2(\pi)}^2.$$

In particular, we can apply this result to the semigroup corresponding to the Langevin diffusion (1.0.1) to obtain a spectral gap criterion for quantitative convergence. However, this result is mainly of use when we are interested in a specific test function f . More generally, it is useful to obtain bounds on the rate of convergence of the law π_t of X_t to the stationary distribution π . Recall from (1.2.16) that the relative density $\rho_t := \pi_t/\pi$ solves the equation $\partial_t \rho_t = \mathcal{L} \rho_t$, i.e., ρ_t is given by $\rho_t = P_t \rho_0$. We can therefore apply the preceding result to $f := \rho_0 - 1$.

For a probability measure μ , we define the **chi-squared divergence**

$$\chi^2(\mu \parallel \pi) := \left\| \frac{d\mu}{d\pi} - 1 \right\|_{L^2(\pi)}^2 = \text{var}_\pi \frac{d\mu}{d\pi} \quad \text{if } \mu \ll \pi,$$

with $\chi^2(\mu \parallel \pi) := \infty$ otherwise. The result can be formulated as follows.

Theorem 1.2.22 (PI and χ^2 decay). *The following are equivalent.*

- 1 *The Markov process satisfies a Poincaré inequality with constant C_{PI} .*
- 2 *For any initial distribution π_0 and all $t \geq 0$,*

$$\chi^2(\pi_t \parallel \pi) \leq \exp\left(-\frac{2t}{C_{\text{PI}}}\right) \chi^2(\pi_0 \parallel \pi).$$

Example 1.2.23. For the Langevin diffusion, the Poincaré inequality reads

$$\text{var}_\pi f \leq C_{\text{PI}} \mathbb{E}_\pi [\|\nabla f\|^2]$$

for all functions $f : \mathbb{R}^d \rightarrow \mathbb{R}$, where $\pi \propto \exp(-V)$.

1.2.3 The Log-Sobolev Inequality and Bakry–Émery Theory

For sampling applications, the convergence result under a Poincaré inequality is not fully satisfactory because the chi-squared divergence at initialization is typically large, scaling exponentially in the dimension. The approach we explore next is to use the **Kullback–Leibler (KL) divergence** $\text{KL}(\cdot \parallel \pi)$ as our objective functional, defined via

$$\text{KL}(\mu \parallel \pi) := \int \frac{d\mu}{d\pi} \log \frac{d\mu}{d\pi} d\pi = \int \log \frac{d\mu}{d\pi} d\mu \quad \text{if } \mu \ll \pi,$$

and $\text{KL}(\mu \parallel \pi) := \infty$ otherwise.

Recall the notation $\rho_t := \pi_t/\pi$ for the relative density of the Markov process w.r.t. π . Since $\partial_t \rho_t = \mathcal{L} \rho_t$, we can calculate via the integration by parts identity that

$$\partial_t \text{KL}(\pi_t \parallel \pi) = \partial_t \int \rho_t \log \rho_t d\pi = \int (\log \rho_t + 1) \mathcal{L} \rho_t d\pi = -\mathcal{E}(\rho_t, \log \rho_t). \quad (1.2.24)$$

Hence, if $\mathcal{E}(\rho_t, \log \rho_t) \gtrsim \text{KL}(\pi_t \parallel \pi)$, then we obtain convergence to equilibrium for the diffusion in KL divergence, at an exponential rate.

Definition 1.2.25. The Markov process is said to satisfy a **log-Sobolev inequality** (LSI) with constant C_{LSI} if for all densities ρ w.r.t. π ,

$$\text{KL}(\rho \pi \parallel \pi) \leq \frac{C_{\text{LSI}}}{2} \mathcal{E}(\rho, \log \rho) .$$

Theorem 1.2.26 (LSI and KL decay). *The following are equivalent.*

- 1 The Markov process satisfies a log-Sobolev inequality with constant C_{LSI} .
- 2 For any initial distribution π_0 and all $t \geq 0$,

$$\text{KL}(\pi_t \parallel \pi) \leq \exp\left(-\frac{2t}{C_{\text{LSI}}}\right) \text{KL}(\pi_0 \parallel \pi) .$$

By *linearizing* the LSI, i.e., by taking $\rho = 1 + \varepsilon f$ for small $\varepsilon > 0$ and expanding both sides of the LSI in powers of ε , one can prove that the LSI implies a PI with constant $C_{\text{PI}} \leq C_{\text{LSI}}$ (Exercise 1.7).

Example 1.2.27. For the Langevin diffusion, the LSI reads

$$\frac{2}{C_{\text{LSI}}} \text{KL}(\mu \parallel \pi) \leq \mathbb{E}_\pi \left\langle \nabla \frac{\mu}{\pi}, \nabla \log \frac{\mu}{\pi} \right\rangle = \mathbb{E}_\mu \left[\left\| \nabla \log \frac{\mu}{\pi} \right\|^2 \right] = 4 \mathbb{E}_\pi \left[\left\| \nabla \sqrt{\frac{\mu}{\pi}} \right\|^2 \right] .$$

The right-hand side of the above expression is important; it is known as the (relative) **Fisher information** $\text{FI}(\mu \parallel \pi) := \mathbb{E}_\mu [\left\| \nabla \log(\mu/\pi) \right\|^2]$. In particular, the Fisher information plays a central role in the study of non-log-concave sampling in Chapter 11.

The LSI often appears in many equivalent forms. For example, another formulation is that for all functions $f : \mathbb{R}^d \rightarrow \mathbb{R}$, it holds that

$$\text{ent}_\pi(f^2) \leq 2C_{\text{LSI}} \mathbb{E}_\pi [\left\| \nabla f \right\|^2] ,$$

where for a function $g : \mathbb{R}^d \rightarrow \mathbb{R}_+$ we define $\text{ent}_\pi(g) := \mathbb{E}_\pi(g \log g) - \mathbb{E}_\pi g \log \mathbb{E}_\pi g$. To verify the equivalence, consider $f = \sqrt{\mu/\pi}$.

Bakry–Émery condition.

Although we have derived two criteria for convergence of the Markov process, namely, the Poincaré inequality and the log-Sobolev inequality, we have not yet addressed when these criteria hold. Introduce the following definition.

Definition 1.2.28. The **iterated carré du champ** is the operator Γ_2 defined via

$$\Gamma_2(f, g) := \frac{1}{2} \{ \mathcal{L}\Gamma(f, g) - \Gamma(f, \mathcal{L}g) - \Gamma(g, \mathcal{L}f) \} .$$

Recalling that $\Gamma(f, g) = \frac{1}{2} \{ \mathcal{L}(fg) - f \mathcal{L}g - g \mathcal{L}f \}$, we see that Γ_2 is defined analogously to Γ , except we replace the bilinear operation of multiplication, $(f, g) \mapsto fg$, by the carré du champ $(f, g) \mapsto \Gamma(f, g)$. Also, similarly to how the carré du champ appears when computing the time derivative of functionals such as the chi-squared divergence and the KL divergence, the iterated carré

du champ appears when computing the *second* time derivative. After some calculations, one arrives at the following criterion.

Definition 1.2.29. The Markov semigroup is said to satisfy the **Bakry–Émery criterion** with constant $\alpha > 0$ if for all functions f ,

$$\Gamma_2(f, f) \geq \alpha \Gamma(f, f).$$

This condition is also known as the **curvature-dimension condition** $\text{CD}(\alpha, \infty)$.

We will prove the following theorem in Chapter 2.

Theorem 1.2.30 (Bakry–Émery). *Consider a diffusion Markov semigroup. Assume that the curvature-dimension condition $\text{CD}(\alpha, \infty)$ holds. Then, a log-Sobolev inequality holds with constant $C_{\text{LSI}} \leq 1/\alpha$.*

We have not explained yet what a *diffusion* Markov semigroup is, but for now we can think of the Langevin diffusion as a fundamental example. The key point is that once the (iterated) carré du champ operators are known, the curvature-dimension condition amounts to an algebraic condition which can be easily checked, which in turn implies the log-Sobolev inequality (and hence the Poincaré inequality by [Exercise 1.7](#)). For the Langevin diffusion, this condition amounts to the following theorem.

Theorem 1.2.31 (CD condition for Langevin). *For the Langevin diffusion (1.0.1), the curvature-dimension condition $\text{CD}(\alpha, \infty)$ holds if and only if the potential V is α -strongly convex.*

Although we have deferred the Markov semigroup proofs of these results to Chapter 2, we will shortly prove these results using the calculus of optimal transport.

Another point to address is the origin of the name “curvature-dimension condition”. In fact this is part of a rich story in which Markov diffusions on Riemannian manifolds capture the geometric features of the ambient space, such as its curvature. A picture emerges in which curvature, concentration, and mixing of the diffusion all intertwine, and only in this context is it appreciated that the curvature-dimension condition is appropriately named. This is discussed more fully in Chapter 2.

1.3 The Geometry of Optimal Transport

In this section, we explain how the space of probability measures equipped with the 2-Wasserstein distance from optimal transport can be formally viewed as a Riemannian manifold. The textbook [Villani \(2003\)](#) is a standard reference for optimal transport; see also [Ambrosio et al. \(2008\)](#); [Villani \(2009b\)](#); [Santambrogio \(2015\)](#); [Chewi et al. \(2025b\)](#) for more detailed treatments of Wasserstein calculus. We remind readers that the “proofs” in this section are only sketched for intuition.

1.3.1 Introduction and Duality Theory

The optimal transport problem can be defined in great generality. Throughout this section, $\mathcal{P}(\mathcal{X})$ denotes the space of probability measures on a space $\mathcal{X}$.

Definition 1.3.1. Let $\mathcal{X}$ and $\mathcal{Y}$ be complete separable metric spaces, and consider a cost functional $c : \mathcal{X} \times \mathcal{Y} \rightarrow [0, \infty]$. The **optimal transport cost** from $\mu \in \mathcal{P}(\mathcal{X})$ to $\nu \in \mathcal{P}(\mathcal{Y})$ with cost c is

$$\mathcal{T}_c(\mu, \nu) := \inf_{\gamma \in \mathcal{C}(\mu, \nu)} \int c(x, y) \gamma(dx, dy), \quad (1.3.2)$$

where $\mathcal{C}(\mu, \nu)$ is the space of *couplings* of (μ, ν) , i.e., the space of probability measures $\gamma \in \mathcal{P}(\mathcal{X} \times \mathcal{Y})$ whose marginals are μ and ν respectively.

A minimizer in this problem is known as an **optimal transport plan**.

An equivalent probabilistic formulation is that $\mathcal{T}_c(\mu, \nu)$ is the infimum of $\mathbb{E} c(X, Y)$ over all pairs of jointly defined random variables (X, Y) such that $X \sim \mu$ and $Y \sim \nu$.

Theorem 1.3.3 (Existence). *If the cost c is lower semicontinuous, then an optimal transport plan always exists.*

Proof One can show that the functional $\gamma \mapsto \int c \, d\gamma$ is lower semicontinuous and that $\mathcal{C}(\mu, \nu)$ is compact, where we use the weak topology on $\mathcal{P}(\mathcal{X} \times \mathcal{Y})$. It is a general fact that lower semicontinuous functions attain their minima on compact sets. $\square$

Historically, optimal transport began with Gaspard Monge who considered the Euclidean cost $c(x, y) := \|x - y\|$ on $\mathbb{R}^d \times \mathbb{R}^d$ (Monge, 1781). Moreover, he considered a slightly different problem in which, rather than searching over all couplings in $\mathcal{C}(\mu, \nu)$, he restricted attention to couplings which are induced by a *mapping* $T : \mathbb{R}^d \rightarrow \mathbb{R}^d$ satisfying $T_{\#}\mu = \nu$; this is known as the **Monge problem**. In the probabilistic interpretation, this corresponds to a pair of random variables $(X, T(X))$ with $X \sim \mu$ and $T(X) \sim \nu$. The physical interpretation of this additional constraint is that no mass from μ be split up before it is transported, which may be reasonable from a modelling perspective but leads to an ill-posed mathematical problem. Indeed, there may not even exist any such mappings T , as is the case when $\mu = \delta_x$ places all of its mass on a single point and ν does not. Consequently, the solution to the Monge problem remained unknown for centuries.

The breakthrough arrived when Leonid Kantorovich formulated the relaxation of the Monge problem introduced in Definition 1.3.1, which is therefore known as the **Kantorovich problem** (Kantorovich, 1942). As the product measure $\mu \otimes \nu$ always belongs to $\mathcal{C}(\mu, \nu)$, we at least know that the constraint set is non-empty, and Theorem 1.3.3 shows that the Kantorovich problem is well-behaved. Moreover, the Kantorovich problem is actually a *convex* problem on $\mathcal{P}(\mathcal{X} \times \mathcal{Y})$; indeed, the objective is linear and the constraint set $\mathcal{C}(\mu, \nu)$ is convex. Hence, one can bring to bear the power of convex duality to study the Kantorovich problem (historically, this study was actually the origin of linear programming).

Although a large part of optimal transport theory can be developed in a general framework as above, for the rest of the section we focus on the case $c(x, y) := \|x - y\|^2$ on $\mathbb{R}^d \times \mathbb{R}^d$ for concreteness.

Definition 1.3.4. The **2-Wasserstein distance** between μ and ν , denoted $W_2(\mu, \nu)$, is defined via

$$W_2^2(\mu, \nu) := \inf_{\gamma \in \mathcal{C}(\mu, \nu)} \int \|x - y\|^2 \gamma(dx, dy). \quad (1.3.5)$$

Write $\mathcal{P}_2(\mathbb{R}^d) := \{\mu \in \mathcal{P}(\mathbb{R}^d) \mid \int \|\cdot\|^2 d\mu < \infty\}$ for the space of probability measures on $\mathbb{R}^d$ with finite second moment.

Duality and optimality.

For this section, it is actually convenient to consider the cost $c(x, y) := \frac{1}{2} \|x - y\|^2$ instead, i.e., we consider $\frac{1}{2} W_2^2(\mu, \nu)$ instead of $W_2^2(\mu, \nu)$.

The key to solving the Kantorovich problem is duality. First observe that the constraint that the first marginal of γ is μ can be written as follows: for every function $f \in L^1(\mu)$, it holds that $\int f(x) \gamma(dx, dy) = \int f(x) \mu(dx)$. Doing the same for the constraint on the second marginal of γ , we can write the Kantorovich problem as an *unconstrained* min-max problem

$$\begin{aligned} \frac{1}{2} W_2^2(\mu, \nu) = \inf_{\gamma \in \mathcal{M}_+(\mathbb{R}^d \times \mathbb{R}^d)} \sup_{\substack{f \in L^1(\mu) \\ g \in L^1(\nu)}} \left\{ \int \frac{\|x - y\|^2}{2} \gamma(dx, dy) + \int f d\mu - \int f(x) \gamma(dx, dy) \right. \\ \left. + \int g d\nu - \int g(y) \gamma(dx, dy) \right\}. \end{aligned}$$

Here, $\mathcal{M}_+(\mathbb{R}^d \times \mathbb{R}^d)$ denotes the space of non-negative finite measures on $\mathbb{R}^d \times \mathbb{R}^d$. Next, if we switch the order of the infimum and the supremum, we arrive at the *dual* optimal transport problem:

$$\begin{aligned} \sup_{\substack{f \in L^1(\mu) \\ g \in L^1(\nu)}} \inf_{\gamma \in \mathcal{M}_+(\mathbb{R}^d \times \mathbb{R}^d)} \left\{ \int f d\mu + \int g d\nu + \int \left[\frac{\|x - y\|^2}{2} - f(x) - g(y) \right] \gamma(dx, dy) \right\} \\ = \sup_{(f, g) \in \mathcal{D}(\mu, \nu)} \left\{ \int f d\mu + \int g d\nu \right\} \end{aligned}$$

where $\mathcal{D}(\mu, \nu)$ is the set of *dual feasible potentials*

$$\mathcal{D}(\mu, \nu) := \left\{ (f, g) \in L^1(\mu) \times L^1(\nu) \mid f(x) + g(y) \leq \frac{\|x - y\|^2}{2} \text{ for } \mu\text{-a.e. } x, \nu\text{-a.e. } y \in \mathbb{R}^d \right\}.$$

Definition 1.3.6. Let $\mu, \nu \in \mathcal{P}_2(\mathbb{R}^d)$. The **dual optimal transport problem** from μ to ν is the optimization problem

$$\sup_{(f, g) \in \mathcal{D}(\mu, \nu)} \left\{ \int f d\mu + \int g d\nu \right\}. \quad (1.3.7)$$

Since $\inf \sup \geq \sup \inf$, the value of the dual problem is always at *most* $\frac{1}{2} W_2^2(\mu, \nu)$. On the other hand, if we find a transport plan γ^* and feasible dual potentials f^*, g^* such that $\frac{1}{2} \int \|x - y\|^2 \gamma^*(dx, dy) = \int f^* d\mu + \int g^* d\nu$, then the primal and dual values coincide and γ^*, f^* , and g^* are all optimal.

By carefully studying the dual problem, we can obtain a wealth of information about the optimal transport problem. Our main goal now is to sketch the following theorem.

Theorem 1.3.8 (Fundamental theorem of optimal transport). *Let $\mu, \nu \in \mathcal{P}_2(\mathbb{R}^d)$. Then, the following assertions hold.*

- 1 (Strong duality) *The value of the dual optimal transport problem from μ to ν equals $\frac{1}{2} W_2^2(\mu, \nu)$.*
- 2 (Existence of optimal dual potentials) *There exists an optimal pair (f^*, g^*) for the dual optimal transport problem.*
- 3 (Characterization of optimality) *The optimal dual potentials are of the form*

$$f^* = \frac{\|\cdot\|^2}{2} - \varphi, \quad g^* = \frac{\|\cdot\|^2}{2} - \varphi^*, \quad (1.3.9)$$

where $\varphi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\infty\}$ is a proper, convex, lower semicontinuous function and φ^* is its convex conjugate. If γ^* denotes the optimal transport plan, then for γ^* -a.e. $(x, y) \in \mathbb{R}^d \times \mathbb{R}^d$, it holds that $\varphi(x) + \varphi^*(y) = \langle x, y \rangle$, i.e., γ^* is supported on the subdifferential of φ .

- 4 (**Brenier's theorem**) *Suppose in addition that μ is absolutely continuous w.r.t. the Lebesgue measure on $\mathbb{R}^d$. Then, the optimal transport plan is unique, and moreover it is induced by an optimal transport map T . The mapping T is characterized as the (μ -almost surely) unique gradient of a proper convex lower semicontinuous function φ which pushes forward μ to ν : $T = \nabla \varphi$ and $(\nabla \varphi)_\# \mu = \nu$.*

Various parts of this theorem can be proven separately; for example, strong duality can be established by rigorously justifying the interchange of infimum and supremum via a high-powered minimax theorem. Instead, we outline a proof which simultaneously establishes all of the above facts.

Outline In the proof, we abbreviate “proper convex lower semicontinuous function” to simply “closed convex function”.

- 1 Optimal transport plans are cyclically monotone. Let γ^* be an optimal transport plan, and suppose that the pairs $(x_1, y_1), \dots, (x_n, y_n)$ lie in the support of γ^* . Then, it should be the case that we cannot “rematch” these points to lower the optimal transport cost, i.e., for every permutation σ of $[n]$ we should have

$$\sum_{i=1}^n \|x_i - y_i\|^2 \leq \sum_{i=1}^n \|x_i - y_{\sigma(i)}\|^2.$$

Equivalently,

$$\sum_{i=1}^n \langle x_i, y_i \rangle \geq \sum_{i=1}^n \langle x_i, y_{\sigma(i)} \rangle. \quad (1.3.10)$$

Indeed, if this condition fails, then it is possible to construct a new transport plan from γ^* by slightly rearranging the mass which has strictly smaller transport cost, which is a contradiction; see [Gangbo and McCann \(1996, Theorem 2.3\)](#).

A subset $S \subseteq \mathbb{R}^d \times \mathbb{R}^d$ is said to be **cyclically monotone** if for all $n \in \mathbb{N}^+$, all pairs $(x_1, y_1), \dots, (x_n, y_n) \in S$, and all permutations σ of $[n]$, the condition (1.3.10) holds. Thus, optimal transport plans are supported on cyclically monotone sets.

- 2 Characterization of cyclically monotone sets. Remarkably, a complete characterization of cyclically monotone sets is known. Suppose φ is convex and differentiable, let $x_1, \dots, x_n \in \mathbb{R}^d$, and let σ be

a permutation of $[n]$. Then, from convexity,

$$\varphi(x_{\sigma^{-1}(i)}) - \varphi(x_i) \geq \langle \nabla \varphi(x_i), x_{\sigma^{-1}(i)} - x_i \rangle. \quad (1.3.11)$$

Summing this over $i \in [n]$, we obtain

$$\sum_{i=1}^n \langle \nabla \varphi(x_i), x_i \rangle \geq \sum_{i=1}^n \langle \nabla \varphi(x_i), x_{\sigma^{-1}(i)} \rangle = \sum_{i=1}^n \langle \nabla \varphi(x_{\sigma(i)}), x_i \rangle.$$

More generally, if φ is not differentiable, then the *subdifferential* of φ at x_i is defined to be the set of vectors $y_i \in \mathbb{R}^d$ such that (1.3.11) holds with y_i replacing $\nabla \varphi(x_i)$. This reasoning shows that the set $\partial \varphi := \{(x, y) \in \mathbb{R}^d \times \mathbb{R}^d \mid y \in \partial \varphi(x)\}$ is a cyclically monotone subset of $\mathbb{R}^d \times \mathbb{R}^d$.

The converse is also true: if $S \subseteq \mathbb{R}^d \times \mathbb{R}^d$ is cyclically monotone, then it is contained in the subdifferential of a closed convex function φ . To prove this, pick any $(x_0, y_0) \in S$ and consider

$$\varphi(x) := \sup \left\{ \sum_{i=0}^n \langle y_i, x_{i+1} - x_i \rangle \mid n \in \mathbb{N}^+, (x_1, y_1), \dots, (x_n, y_n) \in S, x_{n+1} = x \right\}.$$

This characterization is due to [Rockafellar \(1966\)](#).

- 3 **Characterization of dual optimality.** Now that we see the connection between convexity and the primal problem, it is time to do the same for the dual problem. Suppose (f, g) is a feasible dual pair; if we hold f fixed, can we improve g ? The constraint on g says that for all $x, y \in \mathbb{R}^d$,

$$g(y) \leq \frac{\|x - y\|^2}{2} - f(x).$$

Hence, writing $\varphi := \|\cdot\|^2/2 - f$, the optimal choice is

$$g(y) = \inf_{x \in \mathbb{R}^d} \left\{ \frac{\|x - y\|^2}{2} - f(x) \right\} = \frac{\|y\|^2}{2} - \sup_{x \in \mathbb{R}^d} \{ \langle x, y \rangle - \varphi(x) \}.$$

The function φ^* defined by $\varphi^*(y) := \sup_{x \in \mathbb{R}^d} \{ \langle x, y \rangle - \varphi(x) \}$ is known as the *convex conjugate* of φ . To summarize, we have shown that for fixed $f = \|\cdot\|^2/2 - \varphi$, the optimal choice for g is $\|\cdot\|^2/2 - \varphi^*$. Similarly, if we fix $g = \|\cdot\|^2/2 - \varphi^*$, the optimal choice for f is $\|\cdot\|^2/2 - \varphi^{**}$.

We have not yet established existence, but suppose for the moment that an optimal dual pair (f^*, g^*) exists. The preceding reasoning shows that $f^* = \|\cdot\|^2/2 - \varphi$ and $g^* = \|\cdot\|^2/2 - \varphi^*$, where $\varphi^{**} = \varphi$; otherwise the dual pair could be improved. Next, it is known from convex analysis that $\varphi = \varphi^{**}$ if and only if φ is a closed convex function. Thus, optimal dual potentials have the representation (1.3.9).

- 4 **Proof of strong duality.** Now consider the optimal transport plan γ^* (which exists; see [Theorem 1.3.3](#)). We know that γ^* is supported on a cyclically monotone set, which in turn is contained in the subdifferential of a closed convex function φ . Define the functions $f^* := \|\cdot\|^2/2 - \varphi$ and $g^* := \|\cdot\|^2/2 - \varphi^*$; these are dual feasible potentials. Also, it is a standard fact of convex analysis

that $(x, y) \in \partial\varphi$ if and only if $\varphi(x) + \varphi^*(y) = \langle x, y \rangle$. Since the support of γ^* is contained in $\partial\varphi$,

$$\begin{aligned} \frac{1}{2} \int \|x - y\|^2 \gamma^*(dx, dy) &= \int \left(\frac{\|x\|^2}{2} + \frac{\|y\|^2}{2} - \langle x, y \rangle \right) \gamma^*(dx, dy) \\ &= \int \left(\frac{\|x\|^2}{2} + \frac{\|y\|^2}{2} - \varphi(x) - \varphi^*(y) \right) \gamma^*(dx, dy) \\ &= \int \left(\frac{\|\cdot\|^2}{2} - \varphi \right) d\mu + \int \left(\frac{\|\cdot\|^2}{2} - \varphi^* \right) d\nu \\ &= \int f^* d\mu + \int g^* d\nu. \end{aligned}$$

This simultaneously proves that strong duality holds and that (f^*, g^*) is an optimal dual pair.

- 5 For regular measures, optimal transport plans are induced by transport maps. Another fact from convex analysis is that convex functions enjoy some regularity: *a closed convex function φ is differentiable at Lebesgue-a.e. points of the interior of its domain.* Consequently, if μ is absolutely continuous w.r.t. Lebesgue measure, then φ is differentiable μ -a.e. This says that for μ -a.e. $x \in \mathbb{R}^d$, the gradient $\nabla\varphi(x)$ exists and $\partial\varphi(x) = \{\nabla\varphi(x)\}$. Therefore, we can write $\gamma^* = (\text{id}, \nabla\varphi)_\# \mu$. In particular, $(\nabla\varphi)_\# \mu = \nu$, and $\nabla\varphi$ is the optimal transport map from μ to ν .

Our discussion thus far applies to an arbitrary optimal transport plan γ^* . Hence, we have shown that *every optimal transport plan is of the form $(\text{id}, \nabla\varphi)_\# \mu$ for some closed convex function φ .*

- 6 Uniqueness of the optimal transport map. So far, we have not discussed uniqueness of the solution to the Kantorovich problem, and in general uniqueness does not hold. However, in the setting we are currently dealing with (the cost is the squared Euclidean distance and μ is absolutely continuous), we can use additional arguments to establish uniqueness. We show that if $\tilde{\gamma}^* = (\text{id}, \nabla\tilde{\varphi})_\# \mu$ is another optimal transport plan where $\tilde{\varphi}$ is a closed convex function, then $\nabla\varphi = \nabla\tilde{\varphi}$ (μ -a.e.). Note that in particular, it implies that there is only one gradient of a closed convex function which pushes forward μ to ν .

From our above arguments, we see that $(\|\cdot\|^2/2 - \tilde{\varphi}, \|\cdot\|^2/2 - \tilde{\varphi}^*)$ is a dual optimal pair. Therefore,

$$\begin{aligned} \int \{\tilde{\varphi}(x) + \tilde{\varphi}^*(y)\} \gamma^*(dx, dy) &= \int \tilde{\varphi} d\mu + \int \tilde{\varphi}^* d\nu = \int \varphi d\mu + \int \varphi^* d\nu \\ &= \int \{\varphi(x) + \varphi^*(y)\} \gamma^*(dx, dy) = \int \langle x, y \rangle \gamma^*(dx, dy). \end{aligned}$$

Using $\gamma^* = (\text{id}, \nabla\varphi)_\# \mu$, it yields

$$\int \{\tilde{\varphi}(x) + \tilde{\varphi}^*(\nabla\varphi(x)) - \langle x, \nabla\varphi(x) \rangle\} \mu(dx) = 0.$$

On the other hand, by the definition of $\tilde{\varphi}^*$, we have $\tilde{\varphi}(x) + \tilde{\varphi}^*(y) \geq \langle x, y \rangle$ for all $x, y \in \mathbb{R}^d$, with equality if and only if $y \in \partial\tilde{\varphi}(x)$. So, the integrand of the above expression is always non-negative but the integral is zero, which combined with the previous fact shows that $\nabla\varphi(x) \in \partial\tilde{\varphi}(x)$ for μ -a.e. x . But for μ -a.e. x , we also know that $\partial\tilde{\varphi}(x) = \{\nabla\tilde{\varphi}(x)\}$, and we conclude that $\nabla\varphi = \nabla\tilde{\varphi}$ (μ -a.e.). $\square$

We refer to φ as a **Brenier potential**. From convex duality, $\nabla\varphi^* = (\nabla\varphi)^{-1}$. So, if ν is also absolutely continuous, then the optimal transport map from ν to μ is $\nabla\varphi^*$. We often write $T_{\mu \rightarrow \nu} = \nabla\varphi$ for the optimal transport map from μ to ν .

In light of this discussion, it is natural to focus on the following class of measures.

Definition 1.3.12. The space $\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$ is the set of measures in $\mathcal{P}_2(\mathbb{R}^d)$ which are absolutely continuous w.r.t. Lebesgue measure.

Remarks on other costs.

Many of the arguments can be generalized to other costs c . For example, the supports of optimal transport plans can be characterized via c -cyclical monotonicity (generalizing cyclical monotonicity) and optimal dual potentials can be characterized via c -concavity (generalizing convexity). Arguing that the optimal transport plan is induced by a transport map requires additional information about the differentiability of c . See [Exercise 1.10](#).

1.3.2 Riemannian Structure of the Wasserstein Space

Wasserstein space as a metric space.

The following lemma is used to establish that the triangle inequality holds for W_2 .

Lemma 1.3.13 (Gluing lemma). *If $\gamma_{1,2} \in \mathcal{P}(\mathcal{X}_1 \times \mathcal{X}_2)$ and $\gamma_{2,3} \in \mathcal{P}(\mathcal{X}_2 \times \mathcal{X}_3)$ have the same marginal distribution on $\mathcal{X}_2$, then there exists $\gamma \in \mathcal{P}(\mathcal{X}_1 \times \mathcal{X}_2 \times \mathcal{X}_3)$ such that its first two marginals are $\gamma_{1,2}$ and its last two marginals are $\gamma_{2,3}$.*

Proof Let μ denote the common $\mathcal{X}_2$ -marginal of $\gamma_{1,2}$ and $\gamma_{2,3}$. The idea is to first draw $X_2 \sim \mu$. Then, draw X_1 from its conditional distribution given X_2 (according to $\gamma_{1,2}$), and similarly draw X_3 from its conditional distribution given X_2 (according to $\gamma_{2,3}$). Take γ to be the law of the triple (X_1, X_2, X_3) .

The way to formalize this argument is via disintegration of measure. $\square$

Proposition 1.3.14. *The space $(\mathcal{P}_2(\mathbb{R}^d), W_2)$ is a metric space.*

Proof Clearly, W_2 is symmetric in its two arguments. It is also clear that $\mu = \nu$ implies $W_2(\mu, \nu) = 0$. Conversely, if $W_2(\mu, \nu) = 0$, then there exists a coupling (X, Y) of (μ, ν) such that $\|X - Y\|^2 = 0$ a.s., or equivalently $X = Y$ a.s., which gives $\mu = \nu$.

To verify the triangle inequality, we use the gluing lemma. Let $\mu_1, \mu_2, \mu_3 \in \mathcal{P}_2(\mathbb{R}^d)$, let $\gamma_{1,2}^*$ be optimal for (μ_1, μ_2) , and let $\gamma_{2,3}^*$ be optimal for (μ_2, μ_3) . Let γ be obtained by gluing $\gamma_{1,2}^*$ and $\gamma_{2,3}^*$, and let $\gamma_{1,3} \in \mathcal{C}(\mu_1, \mu_3)$ denote the $(1, 3)$ -marginal of γ . Then,

$$\begin{aligned} W_2(\mu_1, \mu_3) &\leq \sqrt{\int \|x_1 - x_3\|^2 \gamma_{1,3}(\mathrm{d}x_1, \mathrm{d}x_3)} \\ &\leq \sqrt{\int \{\|x_1 - x_2\| + \|x_2 - x_3\|\}^2 \gamma(\mathrm{d}x_1, \mathrm{d}x_2, \mathrm{d}x_3)}. \end{aligned}$$

Let $f(x_1, x_2, x_3) := \|x_1 - x_2\|$ and $g(x_1, x_2, x_3) := \|x_2 - x_3\|$. The above expression can be written as

$\|f + g\|_{L^2(\gamma)}$. By applying the triangle inequality in $L^2(\gamma)$,

$$\begin{aligned}
 W_2(\mu_1, \mu_3) &\leq \|f\|_{L^2(\gamma)} + \|g\|_{L^2(\gamma)} \\
 &= \sqrt{\int \|x_1 - x_2\|^2 \gamma(dx_1, dx_2, dx_3)} + \sqrt{\int \|x_2 - x_3\|^2 \gamma(dx_1, dx_2, dx_3)} \\
 &= \sqrt{\int \|x_1 - x_2\|^2 \gamma_{1,2}^*(dx_1, dx_2)} + \sqrt{\int \|x_2 - x_3\|^2 \gamma_{2,3}^*(dx_2, dx_3)} \\
 &= W_2(\mu_1, \mu_2) + W_2(\mu_2, \mu_3) .
 \end{aligned}
 \quad \square$$

Since the next result is technical, we omit the proof.

Proposition 1.3.15. *The metric space $(\mathcal{P}_2(\mathbb{R}^d), W_2)$ is complete and separable. Also, we have $W_2(\mu_n, \mu) \rightarrow 0$ if and only if $\mu_n \rightarrow \mu$ weakly and $\int \|\cdot\|^2 d\mu_n \rightarrow \int \|\cdot\|^2 d\mu$.*

The continuity equation.

Next, we are going to consider dynamics in the space of measures, i.e., curves of measures $t \mapsto \mu_t$. Throughout, we assume these curves are sufficiently nice, in the following sense.

Definition 1.3.16 (Informal). We say that a curve $t \mapsto \mu_t \in \mathcal{P}_{2,ac}(\mathbb{R}^d)$ is **absolutely continuous** if for almost every t ,

$$|\dot{\mu}|(t) := \lim_{s \rightarrow t} \frac{W_2(\mu_s, \mu_t)}{|s - t|} < \infty .$$

The quantity $|\dot{\mu}|$ is called the **metric derivative** of the curve.

More generally, the metric derivative, which represents the magnitude of the velocity of the curve, can be defined on any metric space, whereas an actual derivative (which also requires a notion of direction) cannot; see [Ambrosio et al. \(2008\)](#).

It is helpful to adopt a fluid dynamics analogy in which we think of μ_t as the mass density of a fluid at time t . There are two complementary perspectives on fluid flows: the *Lagrangian* perspective which emphasizes individual particle trajectories, and the *Eulerian* perspective which tracks the evolution of the aggregate fluid density.

Suppose that $X_0 \sim \mu_0$ and that $t \mapsto X_t$ evolves according to the ODE $\dot{X}_t = v_t(X_t)$. Here, $(v_t)_{t \geq 0}$ is a family of vector fields, i.e., mappings $\mathbb{R}^d \rightarrow \mathbb{R}^d$. Since the ODE describes the evolution of the particle trajectory, it is the Lagrangian description of the dynamics. The corresponding Eulerian description is the continuity equation.

Theorem 1.3.17 (Continuity equation). *Let $(v_t)_{t \geq 0}$ be a family of vector fields and suppose that the process $(X_t)_{t \geq 0}$ evolves according to $\dot{X}_t = v_t(X_t)$. Then, the law μ_t of X_t evolves according to the **continuity equation***

$$\partial_t \mu_t + \operatorname{div}(\mu_t v_t) = 0 . \quad (1.3.18)$$

Proof Given a test function $\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\begin{aligned} \int \varphi \partial_t \mu_t &= \partial_t \int \varphi \mu_t = \partial_t \mathbb{E} \varphi(X_t) = \mathbb{E} \langle \nabla \varphi(X_t), \dot{X}_t \rangle = \mathbb{E} \langle \nabla \varphi(X_t), v_t(X_t) \rangle \\ &= \int \langle \nabla \varphi, v_t \rangle \mu_t = - \int \varphi \operatorname{div}(\mu_t v_t). \end{aligned}$$

Since this holds for every φ , we obtain $\partial_t \mu_t + \operatorname{div}(\mu_t v_t) = 0$. $\square$

Remark 1.3.19. We have not justified that the curve $(\mu_t)_{t \geq 0}$ is regular enough to make sense of the PDE (1.3.18) in the classical sense. To be rigorous, we should interpret (1.3.18) in the “weak” sense, that is, in duality with test functions: for all smooth and compactly supported functions φ , $\partial_t \int \varphi d\mu_t = \int \langle \nabla \varphi, v_t \rangle d\mu_t$. For pedagogical reasons, in the subsequent discussion, we simply *assume* that (1.3.18) holds in the classical sense.

Here, Theorem 1.3.17 shows that a family of vector fields and an initial condition μ_0 give rise to a curve of measures. The punchline is that the converse holds: *every* nice curve of measures $(\mu_t)_{t \geq 0}$ can be interpreted as the fluid flow along a family of vector fields, i.e., we can find vector fields $(v_t)_{t \geq 0}$ such that the continuity equation (1.3.18) holds. First, however, note that there is no uniqueness: if $\partial_t \mu_t + \operatorname{div}(\mu_t v_t) = 0$ and for each t , w_t is a vector field satisfying $\operatorname{div}(\mu_t w_t) = 0$, then the continuity equation also holds with the new vector fields $\tilde{v}_t := v_t + w_t$. We will show how to pick a distinguished choice of vector fields $(v_t)_{t \geq 0}$ which can be described in two equivalent ways. First, among all vector fields making the continuity equation hold true, we can choose v_t to minimize $\int \|v_t\|^2 d\mu_t$, which has the physical interpretation of *minimizing kinetic energy*. Second, we can choose v_t to be the gradient of a function; we will see that this is natural in light of the characterization of optimal transport maps.

Theorem 1.3.20 (Curves of measures as fluid flows; cf. Ambrosio et al. (2008, Chapter 8)). *Let $(\mu_t)_{t \geq 0}$ be an absolutely continuous curve of measures.*

- 1 *For any family of vector fields $(\tilde{v}_t)_{t \geq 0}$ such that the continuity equation (1.3.18) holds, we have $|\dot{\mu}|(t) \leq \|\tilde{v}_t\|_{L^2(\mu_t)}$ for almost every t .*
- 2 *Conversely, there exists a unique choice of vector fields $(v_t)_{t \geq 0}$ such that the continuity equation (1.3.18) holds and $\|v_t\|_{L^2(\mu_t)} \leq |\dot{\mu}|(t)$ for almost every t . The choice of vector fields is also characterized by the following property: the continuity equation (1.3.18) holds and for almost every t , v_t belongs to the $L^2(\mu_t)$ closure of gradient vector fields.*

Moreover, the distinguished vector field v_t satisfies, for almost every t ,

$$v_t = \lim_{\delta \searrow 0} \frac{T_{\mu_t \rightarrow \mu_{t+\delta}} - \operatorname{id}}{\delta} \quad \text{in } L^2(\mu_t), \quad (1.3.21)$$

where $T_{\mu_t \rightarrow \mu_{t+\delta}}$ is the optimal transport map from μ_t to $\mu_{t+\delta}$.

Proof sketch 1 Proof of the first statement. Let $\delta > 0$ and consider the flow map $F_{t,t+\delta}$ defined as follows. Given any initial point $x_t \in \mathbb{R}^d$, consider the ODE $\dot{x}_t = \tilde{v}_t(x_t)$ started at x_t . Then, $F_{t,t+\delta}$ maps x_t to the solution $x_{t+\delta}$ of the ODE at time $t + \delta$.

If $X_t \sim \mu_t$, then the continuity equation implies $F_{t,t+\delta}(X_t) \sim \mu_{t+\delta}$, i.e., $F_{t,t+\delta}$ is a valid transport

map from μ_t to $\mu_{t+\delta}$. Hence, we can estimate

$$\frac{W_2(\mu_t, \mu_{t+\delta})}{\delta} \leq \sqrt{\int \frac{\|F_{t,t+\delta} - \text{id}\|^2}{\delta^2} d\mu_t}.$$

However, $F_{t,t+\delta} - \text{id} = \delta \tilde{v}_t + o(\delta)$, so letting $\delta \searrow 0$ we obtain $|\dot{\mu}|(t) \leq \|\tilde{v}_t\|_{L^2(\mu_t)}$. (Actually, to prove this statement we should also consider the limit $\delta \nearrow 0$ for negative δ , but it is clear that the same argument works.)

- 2 Uniqueness of the optimal vector field. Suppose we find $(v_t)_{t \geq 0}$ satisfying the continuity equation and such that $\|v_t\|_{L^2(\mu_t)} \leq |\dot{\mu}|(t)$. In light of the first statement, it implies that the zero vector field is the minimizer of $\|v_t + w_t\|_{L^2(\mu_t)}$ among all vector fields w_t such that $\text{div}(\mu_t w_t) = 0$. This is a strictly convex problem so the minimizer is unique, meaning that $(v_t)_{t \geq 0}$ is uniquely determined.
- 3 Gradient vector fields are optimal. Here, we show that if the continuity equation holds for the family of vector fields $(v_t)_{t \geq 0}$ and that $v_t = \nabla \psi_t$ for all t , then the vector fields are optimal.

There are at least two ways of seeing why gradient vector fields should be optimal. First, the continuity equation is equivalent to requiring that for all test functions $\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$, it holds that $\partial_t \int \varphi d\mu_t = \int \langle \nabla \varphi, v_t \rangle d\mu_t$. In this expression, the vector field v_t only enters through inner products with gradients. To put it another way, if we consider the space $S := \{\nabla \psi \mid \psi : \mathbb{R}^d \rightarrow \mathbb{R}\}$ of gradients, viewed as a subspace of $L^2(\mu_t)$, then we can write $L^2(\mu_t) = S \oplus S^\perp$ (actually, to make this valid we should take the closure of S , but we will ignore this detail). If we decompose v_t according to this direct sum, then $v_t = \nabla \psi_t + w_t$ for some function ψ_t and some w_t which is orthogonal (in $L^2(\mu_t)$) to S . If we replace v_t by $\nabla \psi_t$, then the continuity equation continues to hold, but we have only made the norm $\|v_t\|_{L^2(\mu_t)}$ smaller, hence the optimal choice of v_t should lie in S .

The second line of reasoning comes from the proof of the first statement: the reason why the metric derivative $|\dot{\mu}|(t)$ was upper bounded by $\|\tilde{v}_t\|_{L^2(\mu_t)}$ is because the flow map corresponding to $(\tilde{v}_t)_{t \geq 0}$ furnishes a possibly suboptimal transport map. To fix this, the flow map for the optimal $(v_t)_{t \geq 0}$ should be approximately equal to the *optimal* transport map, i.e., $v_t \approx (T_{\mu_t \rightarrow \mu_{t+\delta}} - \text{id})/\delta$. From the fundamental theorem of optimal transport, however, $T_{\mu_t \rightarrow \mu_{t+\delta}}$ is the gradient of a convex function, so in the limit v_t should be a gradient field as well.

Instead of using these arguments, we will instead provide a proof based on direct computation. If $v_t = \nabla \psi_t$, the continuity equation shows that

$$\begin{aligned} \frac{\int \psi_t d\mu_{t+\delta} - \int \psi_t d\mu_t}{\delta} &= \int \psi_t \partial_t \mu_t + o(1) = - \int \psi_t \text{div}(\mu_t \nabla \psi_t) + o(1) \\ &= \int \|\nabla \psi_t\|^2 d\mu_t + o(1). \end{aligned}$$

On the other hand,

$$\begin{aligned} \frac{\int \psi_t d\mu_{t+\delta} - \int \psi_t d\mu_t}{\delta} &= \int \frac{\psi_t \circ T_{\mu_t \rightarrow \mu_{t+\delta}} - \psi_t}{\delta} d\mu_t \\ &= \int \left\langle \nabla \psi_t, \frac{T_{\mu_t \rightarrow \mu_{t+\delta}} - \text{id}}{\delta} \right\rangle d\mu_t + o(1) \\ &\leq \sqrt{\int \|\nabla \psi_t\|^2 d\mu_t} \frac{W_2(\mu_t, \mu_{t+\delta})}{\delta} + o(1). \end{aligned}$$

Taking $\delta \searrow 0$ yields $\|v_t\|_{L^2(\mu_t)} = \|\nabla \psi_t\|_{L^2(\mu_t)} \leq |\dot{\mu}|(t)$.

- 4 Existence of optimal vector fields. Finally, one can show for instance that vector fields defined via limits of transport maps as in (1.3.21) indeed satisfy the continuity equation and are gradient vector fields, and are therefore optimal. The details are omitted. $\square$

From the theorem, we learn that the optimal vector field v_t satisfies $\|v_t\|_{L^2(\mu_t)} = |\dot{\mu}|(t)$. On the other hand, the metric derivative is supposed to be the “magnitude of the velocity”. Our next goal is to interpret v_t as the velocity vector to the curve, and $\|v_t\|_{L^2(\mu_t)}$ as its norm, all through the lens of Riemannian geometry.

Background on Riemannian geometry.

In the spirit of informality, we give a description of what a Riemannian manifold entails, rather than a precise definition. A **manifold** $\mathcal{M}$ is a space which is locally homeomorphic to a Euclidean space. At each point $p \in \mathcal{M}$, there is an associated vector space $T_p\mathcal{M}$, called the **tangent space** at p , which is the space of all possible velocities of curves passing through p . The whole structure should be smooth: the tangent spaces should vary smoothly in a suitable sense.

A **Riemannian metric** is a smoothly varying choice of inner products $p \mapsto \langle \cdot, \cdot \rangle_p$ on the tangent spaces. The metric allows us to, e.g., locally measure the angles between two intersecting curves. For our purposes, it is important to note that the metric allows us to define the steepest descent direction for an objective function, which in turn allows us to consider gradient flows.

The Riemannian metric induces a distance function (in the sense of metric spaces) via

$$d(p, q) := \inf \left\{ \int_0^1 \|\dot{\gamma}(t)\|_{\gamma(t)} dt \mid \gamma : [0, 1] \rightarrow \mathcal{M}, \gamma(0) = p, \gamma(1) = q \right\}. \quad (1.3.22)$$

Here, $\dot{\gamma}(t)$ denotes the tangent vector to the curve at time t . Note that the norm of the tangent vector is measured w.r.t. the inner product on the tangent space $T_{\gamma(t)}\mathcal{M}$, hence we write $\|\dot{\gamma}(t)\|_{\gamma(t)}$. If the infimum is achieved by a curve γ , then γ is referred to as a **geodesic** (a shortest path); if $t \mapsto \|\dot{\gamma}(t)\|_{\gamma(t)}$ is constant, then it is called a **constant-speed geodesic**. From now on, we will only consider constant-speed geodesics, and the words “constant speed” will be dropped for brevity.

Given a functional $\mathcal{F} : \mathcal{M} \rightarrow \mathbb{R}$, the gradient of $\mathcal{F}$ at p is defined to be the unique element $\nabla \mathcal{F}(p) \in T_p\mathcal{M}$ such that for all curves $(p_t)_{t \in \mathbb{R}}$ passing through p at time 0 with velocity $v \in T_p\mathcal{M}$, it holds that $\partial_t|_{t=0} \mathcal{F}(p_t) = \langle \nabla \mathcal{F}(p), v \rangle_p$.

Wasserstein space as a Riemannian manifold.

Based on our discussion thus far, it is natural to define the tangent space at $\mu \in \mathcal{P}_{2,ac}(\mathbb{R}^d)$ as

$$T_\mu \mathcal{P}_{2,ac}(\mathbb{R}^d) := \overline{\{\nabla \psi \mid \psi \in C_c^\infty(\mathbb{R}^d)\}}^{L^2(\mu)},$$

where the notation denotes taking the $L^2(\mu)$ closure. Equivalently,

$$T_\mu \mathcal{P}_{2,ac}(\mathbb{R}^d) = \overline{\{\lambda(T - \text{id}) \mid \lambda > 0, T \text{ is an optimal transport map}\}}^{L^2(\mu)}.$$

We equip the tangent space $T_\mu \mathcal{P}_{2,ac}(\mathbb{R}^d)$ with the $L^2(\mu)$ norm, which gives a Riemannian metric. This does not define a genuine Riemannian manifold (e.g., it is not locally homeomorphic to a Euclidean space or even a Hilbert space), but we treat it as one for the purpose of developing calculation rules.

If the continuity equation $\partial_t \mu_t + \text{div}(\mu_t v_t) = 0$ holds and $v_t \in T_{\mu_t} \mathcal{P}_{2,ac}(\mathbb{R}^d)$, then v_t is the tangent vector to the curve at time t . The condition $v_t \in T_{\mu_t} \mathcal{P}_{2,ac}(\mathbb{R}^d)$ is equivalent to saying that v_t is the optimal vector field considered in Theorem 1.3.20.

There are two questions to address. First, is this Riemannian structure compatible with the 2-Wasserstein distance? In other words, we know that a Riemannian metric induces a distance function; is the distance function induced by the Riemannian structure of $\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$ equal to W_2 ? Second, what are the geodesics? We answer both questions via the following theorem.

Theorem 1.3.23 (Wasserstein geodesics). *Let $\mu_0, \mu_1 \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$. Then,*

$$W_2(\mu_0, \mu_1) = \inf \left\{ \int_0^1 \|v_t\|_{L^2(\mu_t)} dt \mid \partial_t \mu_t + \text{div}(\mu_t v_t) = 0 \right\}. \quad (1.3.24)$$

The infimum is achieved as follows. Let $X_0 \sim \mu_0$ and $X_1 \sim \mu_1$ be optimally coupled random variables, let $X_t := (1-t)X_0 + tX_1$, and let $\mu_t := \text{law}(X_t)$. Then, $(\mu_t)_{t \in [0,1]}$ is the unique constant-speed geodesic joining μ_0 to μ_1 .

Proof Suppose that $\partial_t \mu_t + \text{div}(\mu_t v_t) = 0$. Then, $\int_0^1 \|v_t\|_{L^2(\mu_t)} dt \geq \int_0^1 |\dot{\mu}|(t) dt$. For a partition $0 = t_0 < t_1 < \dots < t_k = 1$,

$$W_2(\mu_0, \mu_1) \leq \sum_{i=1}^k W_2(\mu_{t_{i-1}}, \mu_{t_i}) = \sum_{i=1}^k \frac{W_2(\mu_{t_{i-1}}, \mu_{t_i})}{t_i - t_{i-1}} (t_i - t_{i-1}).$$

As the size of the partition tends to zero, we obtain $W_2(\mu_0, \mu_1) \leq \int_0^1 |\dot{\mu}|(t) dt$. This shows that $W_2(\mu_0, \mu_1)$ is at most the value of the infimum.

To show that equality holds, let X_t be defined as in the theorem statement and note that $\mathbb{E}[\|\dot{X}_t\|^2] = \|v_t\|_{L^2(\mu_t)}^2$ by the correspondence of the Lagrangian and Eulerian perspectives. (This can be verified by writing the vector field explicitly as $v_t = (T_1 - \text{id}) \circ T_t^{-1}$, where $T_t := (1-t)\text{id} + tT_{\mu_0 \rightarrow \mu_1}$ —exercise!) Since $\mathbb{E}[\|\dot{X}_t\|^2] = \mathbb{E}[\|X_1 - X_0\|^2] = W_2^2(\mu_0, \mu_1)$ does not depend on time, the curve has constant speed, and $\int_0^1 \|v_t\|_{L^2(\mu_t)} dt = W_2(\mu_0, \mu_1)$.

To show uniqueness, again work in the Lagrangian perspective: suppose we have a process $(X_t)_{t \geq 0}$ such that $t \mapsto \mathbb{E}[\|\dot{X}_t\|^2]$ is constant, and $X_0 \sim \mu_0, X_1 \sim \mu_1$. Then, we have

$$W_2^2(\mu_0, \mu_1) \leq \mathbb{E}[\|X_1 - X_0\|^2] = \mathbb{E}\left[\left\|\int_0^1 \dot{X}_t dt\right\|^2\right] \leq \mathbb{E} \int_0^1 \|\dot{X}_t\|^2 dt = \left(\int_0^1 \|v_t\|_{L^2(\mu_t)} dt\right)^2$$

where the last equality follows from the constant-speed assumption. In order for the first inequality to be equality, (X_0, X_1) is an optimal coupling. In order for the second inequality to be equality, strict convexity of $\|\cdot\|^2$ implies that $\dot{X}_t$ is constant in time and equal to its average $\int_0^1 \dot{X}_t dt = X_1 - X_0$. $\square$

Definition 1.3.25. Let $\mu_0, \mu_1 \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$, and let $X_0 \sim \mu_0, X_1 \sim \mu_1$ be optimally coupled. Let $X_t := (1-t)X_0 + tX_1$, and let $\mu_t := \text{law}(X_t)$. Then, the curve $(\mu_t)_{t \in [0,1]}$ is called the **Wasserstein geodesic** joining μ_0 to μ_1 . It is also called the **displacement interpolation** or **McCann's interpolation**.

Exponential map and logarithmic map.

On a Riemannian manifold $\mathcal{M}$, the Riemannian **exponential map** $\exp_p : T_p \mathcal{M} \rightarrow \mathcal{M}$ takes a tangent vector $v \in T_p \mathcal{M}$ to the endpoint at time 1 of the constant-speed geodesic emanating from p with velocity v . The **logarithmic map** is then defined to be the inverse mapping $\log_p : \mathcal{M} \rightarrow T_p \mathcal{M}$.

Actually, in general, the exponential map is only defined on a subset of the tangent space, because in many manifolds (e.g., the sphere), geodesics cannot continue indefinitely while remaining shortest paths between their endpoints. On Euclidean space $\mathbb{R}^d$, we have $\exp_p(v) = p + v$ and $\log_p(q) = q - p$.

We can identify these maps for $(\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d), W_2)$. If $(\mu_t)_{t \in [0,1]}$ is a Wasserstein geodesic and $\nabla \varphi_{\mu_0 \rightarrow \mu_1}$ is the optimal transport map from μ_0 to μ_1 , then the tangent vector to the geodesic at time 0 is $\nabla \varphi_{\mu_0 \rightarrow \mu_1} - \text{id}$. This implies that $\log_{\mu_0}(\mu_1) = \nabla \varphi_{\mu_0 \rightarrow \mu_1} - \text{id}$. The inverse mapping is then given as follows: if $\nabla \psi \in T_{\mu_0} \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$ is such that $\text{id} + \nabla \psi$ is the gradient of a convex function, then $\exp_{\mu_0}(\nabla \psi) = (\text{id} + \nabla \psi)_\# \mu_0$.

Geodesically convex functionals.

Over a Riemannian manifold $\mathcal{M}$, the correct way to define convexity is as follows.

Definition 1.3.26. Let $\mathcal{M}$ be a Riemannian manifold and let $\mathcal{F} : \mathcal{M} \rightarrow \mathbb{R}$ be smooth. For $\alpha \in \mathbb{R}$, we say that $\mathcal{F}$ is **α -geodesically convex** if one of the following equivalent conditions holds:

- 1 For all geodesics $(p_t)_{t \in [0,1]}$ and $t \in [0, 1]$,

$$\mathcal{F}(p_t) \leq (1-t) \mathcal{F}(p_0) + t \mathcal{F}(p_1) - \frac{\alpha t(1-t)}{2} d(p_0, p_1)^2,$$

where d is the induced Riemannian distance (1.3.22).

- 2 For all $p, q \in \mathcal{M}$,

$$\mathcal{F}(q) \geq \mathcal{F}(p) + \langle \nabla \mathcal{F}(p), \log_p(q) \rangle_p + \frac{\alpha}{2} d(p, q)^2.$$

Here, ∇ denotes the Riemannian gradient.

- 3 For all constant-speed geodesics $(p_t)_{t \in [0,1]}$ with tangent vector $v_0 \in T_{p_0} \mathcal{M}$ at time 0,

$$\partial_t^2|_{t=0} \mathcal{F}(p_t) \geq \alpha \|v_0\|_{p_0}^2.$$

1.4 The Langevin SDE as a Wasserstein Gradient Flow

We are now ready to interpret the Langevin diffusion (1.0.1) as a gradient flow in the Wasserstein space $(\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d), W_2)$. Once we have done so, we will quickly deduce convergence results for the Langevin diffusion based on gradient flow computations.

1.4.1 Derivation of the Gradient Flow

Let $\mathcal{F} : \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d) \rightarrow \mathbb{R} \cup \{\infty\}$ be a functional over the Wasserstein space. We now compute the Wasserstein gradient of $\mathcal{F}$ at μ , i.e., the element $\nabla_{W_2} \mathcal{F}(\mu) \in T_\mu \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$ such that for every curve $(\mu_t)_{t \geq 0}$ with $\mu_0 = \mu$, if v_0 is the tangent vector to the curve at time 0, then $\partial_t|_{t=0} \mathcal{F}(\mu_t) = \langle \nabla_{W_2} \mathcal{F}(\mu), v_0 \rangle_\mu$, where $\langle \cdot, \cdot \rangle_\mu$ is the inner product on $T_\mu \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$.

We will give a formula in terms of the **first variation** of $\mathcal{F}$ at μ , denoted $\delta \mathcal{F}(\mu)$. The first variation is a function $\mathbb{R}^d \rightarrow \mathbb{R}$ which satisfies $\partial_t|_{t=0} \mathcal{F}(\mu_t) = \int \delta \mathcal{F}(\mu) \partial_t|_{t=0} \mu_t$. This is almost the same as the $L^2(\mathbf{m})$ gradient of $\mathcal{F}$, where $\mathbf{m}$ is the Lebesgue measure on $\mathbb{R}^d$, except for a few differences: (1) there is no guarantee that $\delta \mathcal{F}(\mu) \in L^2(\mathbf{m})$; (2) in order to consider the $L^2(\mathbf{m})$ gradient, we would want $\mathcal{F}$ to be a functional defined over all of $L^2(\mathbf{m})$, not just probability densities, and similarly we would have to consider all curves in $L^2(\mathbf{m})$ rather than curves of probability densities.

As a consequence of looking only at probability densities, the first variation is only defined up to an additive constant. Indeed, $\partial_t|_{t=0}\mu_t$ always integrates to 0, so we can add any constant to the first variation. This does not cause any ambiguity, as we now see.

Recall that v_t is the tangent vector to the curve of measures at time t if the continuity equation $\partial_t\mu_t + \operatorname{div}(\mu_t v_t) = 0$ holds and $v_t \in T_{\mu_t}\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$. Using the continuity equation with a curve such that $\mu_0 = \mu$,

$$\partial_t|_{t=0}\mathcal{F}(\mu_t) = \int \delta\mathcal{F}(\mu) \partial_t|_{t=0}\mu_t = - \int \delta\mathcal{F}(\mu) \operatorname{div}(v_0\mu) = \int \langle \nabla\delta\mathcal{F}(\mu), v_0 \rangle d\mu.$$

(Here, the ∇ is the Euclidean gradient.) Since $\nabla\delta\mathcal{F}(\mu)$ is the gradient of a function, from our characterization of the tangent space we know that $\nabla\delta\mathcal{F}(\mu) \in T_\mu\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$. Therefore, the equation above says that the Wasserstein gradient of $\mathcal{F}$ at μ is $\nabla\delta\mathcal{F}(\mu)$.

Theorem 1.4.1 (Wasserstein gradient). *Let $\mathcal{F} : \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d) \rightarrow \mathbb{R} \cup \{\infty\}$ be a functional. Suppose that a first variation $\delta\mathcal{F}$ exists, and that $\nabla\delta\mathcal{F}(\mu) \in T_\mu\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$. Then,*

$$\nabla_{w_2}\mathcal{F}(\mu) = \nabla\delta\mathcal{F}(\mu).$$

Since we take the Euclidean gradient of the first variation, the fact that the first variation is only defined up to additive constant does not bother us.

The **Wasserstein gradient flow** of $\mathcal{F}$ is by definition a curve of measures $(\mu_t)_{t \geq 0}$ such that its tangent vector v_t at time t is $v_t = -\nabla_{w_2}\mathcal{F}(\mu_t)$. Substituting this into the continuity equation (1.3.18), we obtain the gradient flow equation

$$\partial_t\mu_t = \operatorname{div}(\mu_t \nabla_{w_2}\mathcal{F}(\mu_t)) = \operatorname{div}(\mu_t \nabla\delta\mathcal{F}(\mu_t)).$$

Example 1.4.2. Consider $\mathcal{F} = \operatorname{KL}(\cdot \| \pi)$ where $\pi \propto \exp(-V)$. This functional can be rewritten as

$$\mathcal{F}(\mu) = \int \mu \log \frac{\mu}{\pi} = \int V d\mu + \int \mu \log \mu + \text{constant}.$$

These first two terms have the interpretation of *energy* and (negative) *entropy*. From this, we can compute that

$$\delta\mathcal{F}(\mu) = V + \log \mu + \text{constant}$$

and therefore

$$\nabla_{w_2}\mathcal{F}(\mu) = \nabla V + \nabla \log \mu = \nabla \log \frac{\mu}{\pi}.$$

The Wasserstein gradient flow of $\mathcal{F}$ satisfies

$$\partial_t\mu_t = \operatorname{div}(\mu_t \nabla \log \frac{\mu_t}{\pi}).$$

Comparing with the Fokker–Planck equation $\partial_t\pi_t = \mathcal{L}^*\pi_t$ and the form of the adjoint generator $\mathcal{L}^*$ for the Langevin diffusion (see Example 1.2.8), we obtain a truly remarkable fact: *the law $(\pi_t)_{t \geq 0}$ of the Langevin diffusion with potential V is the Wasserstein gradient flow of $\operatorname{KL}(\cdot \| \pi)$.*

The calculus of optimal transport was introduced by Otto in Otto (2001), and it is often known as

Otto calculus; the interpretation of the Langevin diffusion in this context was put forth in the seminal work [Jordan et al. \(1998\)](#). The paper [Otto \(2001\)](#) also raises and answers a salient question: given that we can view dynamics as gradient flows in different ways (e.g. the Langevin diffusion can be either viewed as the gradient flow of the Dirichlet energy in $L^2(\pi)$, or the gradient flow of the KL divergence in $\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$), what makes us prefer one gradient flow structure over another? Otto argues that the Wasserstein geometry is particularly natural because it cleanly separates out two aspects of the problem: the *geometry* of the ambient space, which is reflected in the metric on $\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$, and the *objective functional*. Moreover, the objective functional in the Wasserstein perspective is physically intuitive because it has an interpretation in *thermodynamics*. From a sampling standpoint, the Wasserstein geometry is undoubtedly more compelling and useful.

In our exposition, we focused on the calculation rules for Wasserstein gradient flows, but this is not how they are normally defined. Instead, one usually considers a sequence of discrete approximations to the gradient flow and proves that there is a limiting curve; this is called the minimizing movements scheme and it is developed in detail in [Ambrosio et al. \(2008\)](#).

1.4.2 Convexity of the KL Divergence

The key to studying gradient flows is to understand the convexity properties of the objective functional. For the specific functional $\mathcal{F} := \text{KL}(\cdot \parallel \pi)$ with target $\pi = \exp(-V)$, our next goal is therefore to compute the Wasserstein Hessian of $\mathcal{F}$. When we computed Wasserstein gradients, we were free to differentiate $\mathcal{F}$ along any curve $(\mu_t)_{t \geq 0}$ of measures, but we have to be more careful when computing the Hessian. If we take a function $f : \mathbb{R}^d \rightarrow \mathbb{R}$ on Euclidean space and a curve $(x_t)_{t \geq 0}$, then $\partial_t f(x_t) = \langle \nabla f(x_t), \dot{x}_t \rangle$ and

$$\partial_t^2 f(x_t) = \langle \nabla^2 f(x_t) \dot{x}_t, \dot{x}_t \rangle + \langle \nabla f(x_t), \ddot{x}_t \rangle.$$

Here, instead of just obtaining the Hessian, we have an additional term. However, if $t \mapsto x_t$ is a constant-speed *geodesic*, then it has no acceleration ($\ddot{x}_t = 0$), and the extra term vanishes.

In the same way, let $(\mu_t)_{t \in [0,1]}$ denote a Wasserstein geodesic. Explicitly, if T denotes the optimal transport map from μ_0 to μ_1 , then $\mu_t = [(1-t)\text{id} + tT]_{\#}\mu_0$. We will calculate $\partial_t^2|_{t=0}\mathcal{F}(\mu_t)$, as a function of the tangent vector $T - \text{id} \in T_{\mu_0}\mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$; this is interpreted as $\langle \nabla_{W_2}^2 \mathcal{F}(\mu_0) (T - \text{id}), T - \text{id} \rangle_{\mu_0}$. If we can lower bound this by $\alpha \|T - \text{id}\|_{\mu_0}^2$, for all μ_0 and all optimal transport maps T , it means that $\mathcal{F}$ is α -strongly convex.

Write $\mathcal{E}(\mu) := \int V d\mu$ for the energy and $\mathcal{H}(\mu) := \int \mu \log \mu$ for the entropy. We deal with the two terms separately. First, for $X_t = (1-t)X_0 + tT(X_0)$ and $X_0 \sim \mu_0$,

$$\begin{aligned} \partial_t \mathcal{E}(\mu_t) &= \partial_t \mathbb{E} V(X_t) = \mathbb{E} \langle \nabla V(X_t), \dot{X}_t \rangle = \mathbb{E} \langle \nabla V(X_t), T(X_0) - X_0 \rangle, \\ \partial_t^2|_{t=0} \mathcal{E}(\mu_t) &= \mathbb{E} \langle \nabla^2 V(X_0) (T(X_0) - X_0), T(X_0) - X_0 \rangle. \end{aligned}$$

If V is α -strongly convex, then this is lower bounded by $\alpha \|T - \text{id}\|_{\mu_0}^2$, so that $\mathcal{E}$ is α -strongly convex.

The entropy is slightly trickier. Write $T_t := (1-t)\text{id} + tT$. Since $(T_t)_{\#}\mu_0 = \mu_t$, the change of variables formula shows that

$$\frac{\mu_0}{\mu_t \circ T_t} = \det \nabla T_t.$$

Therefore,

$$\begin{aligned}\mathcal{H}(\mu_t) &= \int \mu_t \log \mu_t = \int \mu_0 \log(\mu_t \circ T_t) = \int \mu_0 \log \frac{\mu_0}{\det \nabla T_t} \\ &= \mathcal{H}(\mu_0) - \int \mu_0 \log \det((1-t)I_d + t \nabla T) .\end{aligned}\quad (1.4.3)$$

Already from the fact that $-\log \det$ is convex on the space of positive semidefinite matrices, we can see that $\partial_t^2 \mathcal{H}(\mu_t) \geq 0$. A more careful computation ([Exercise 1.15](#)) based on the derivatives of $-\log \det$ shows that

$$\partial_t^2|_{t=0} \mathcal{H}(\mu_t) = \int \|\nabla T - I_d\|_{\text{HS}}^2 d\mu_0 \geq 0 . \quad (1.4.4)$$

We have obtained the following result.

Theorem 1.4.5 (Convexity of the KL divergence). *If $\pi \propto \exp(-V)$, where V is α -strongly convex, then $\text{KL}(\cdot \parallel \pi)$ is also α -strongly convex along Wasserstein geodesics.*

The converse is also true: if $\text{KL}(\cdot \parallel \pi)$ is α -strongly convex along Wasserstein geodesics, then necessarily V is α -strongly convex.

Consequences of strong convexity.

First, we describe consequences of strong convexity for gradient flows in general. Let $(\mu_t)_{t \geq 0}, (\nu_t)_{t \geq 0}$ denote gradient flows for a functional $\mathcal{F}$ which is minimized at $\mu_\star$. More generally, the discussion here applies to gradient flows over Riemannian manifolds, with the appropriate change of definitions. We refer to [Chewi \(2025\)](#) for a reference on the corresponding results for Euclidean gradient flows.

We begin with a preliminary fact.

Theorem 1.4.6 (Derivative of W_2). *For $v \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$, the Wasserstein gradient of $\mu \mapsto W_2^2(\mu, v)$ at μ is given by $-2(T_{\mu \rightarrow v} - \text{id})$.*

Proof See [Villani \(2009b\)](#), Theorem 23.9. □

In general, on a Riemannian manifold, the gradient of $d(\cdot, q)^2$ at p is $-2 \log_p(q)$. Check that this formula makes sense on Euclidean space $\mathbb{R}^d$.

First, suppose that $\mathcal{F}$ is α -convex along Wasserstein geodesics; we allow $\alpha < 0$. Recall from [Definition 1.3.26](#) that this is equivalent to the first-order condition

$$\mathcal{F}(v) \geq \mathcal{F}(\mu) + \langle \nabla_{W_2} \mathcal{F}(\mu), T_{\mu \rightarrow v} - \text{id} \rangle_\mu + \frac{\alpha}{2} W_2^2(\mu, v) \quad \text{for all } \mu, v \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d) . \quad (1.4.7)$$

We now justify the following facts.

- The gradient flow is contractive:³ for all $t \geq 0$,

$$W_2(\mu_t, \nu_t) \leq \exp(-\alpha t) W_2(\mu_0, \nu_0) .$$

See [Exercise 1.17](#). In fact, contractivity is equivalent to α -convexity. When $\nu_0 = \mu_\star$, then $\nu_t = \mu_\star$ for all $t \geq 0$, so this also yields convergence to the minimizer.

³“Contractive” is a misnomer when $\alpha < 0$.

- The gradient flow converges in function value: for all $t \geq 0$,

$$\mathcal{F}(\mu_t) - \inf \mathcal{F} \leq \frac{\alpha}{2(\exp(\alpha t) - 1)} W_2^2(\mu_0, \mu_\star). \quad (1.4.8)$$

In particular, for $\alpha = 0$, it yields

$$\mathcal{F}(\mu_t) - \inf \mathcal{F} \leq \frac{1}{2t} W_2^2(\mu_0, \mu_\star).$$

To justify this, we use [Theorem 1.4.6](#) and (1.4.7) to obtain

$$\partial_t W_2^2(\mu_t, \mu_\star) = 2 \langle \nabla_{W_2} \mathcal{F}(\mu_t), T_{\mu_t \rightarrow \mu_\star} - \text{id} \rangle_{\mu_t} \leq -2(\mathcal{F}(\mu_t) - \inf \mathcal{F}) - \alpha W_2^2(\mu_t, \mu_\star).$$

By differentiating $t \mapsto \exp(\alpha t) W_2^2(\mu_t, \mu_\star)$, one readily obtains (1.4.8).

Convexity is not necessary for the gradient flow to converge. It is always true that

$$\partial_t (\mathcal{F}(\mu_t) - \inf \mathcal{F}) = \langle \nabla_{W_2} \mathcal{F}(\mu_t), \underbrace{-\nabla_{W_2} \mathcal{F}(\mu_t)}_{\text{tangent vector of the curve}} \rangle_{\mu_t} = -\|\nabla_{W_2} \mathcal{F}(\mu_t)\|_{\mu_t}^2. \quad (1.4.9)$$

Thus, $t \mapsto \mathcal{F}(\mu_t)$ is monotonically decreasing, and lower bounds on the gradient norm yield rates of convergence. One such condition, known as a *gradient domination* condition, or a **Polyak–Łojasiewicz (PL) inequality**, reads

$$\|\nabla_{W_2} \mathcal{F}(\mu)\|_{\mu}^2 \geq 2\alpha (\mathcal{F}(\mu) - \inf \mathcal{F}) \quad \text{for all } \mu \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d). \quad (1.4.10)$$

- The functional $\mathcal{F}$ satisfies (1.4.10) if and only if

$$\mathcal{F}(\mu_t) - \inf \mathcal{F} \leq \exp(-2\alpha t) (\mathcal{F}(\mu_0) - \inf \mathcal{F})$$

for all gradient flows $(\mu_t)_{t \geq 0}$ for $\mathcal{F}$. This is a consequence of Grönwall's lemma ([Lemma 1.2.20](#)).

- The condition (1.4.10) is implied by strong convexity: starting from (1.4.7), we take $\nu = \mu_\star$, yielding

$$\mathcal{F}(\mu) - \inf \mathcal{F} \leq -\langle \nabla_{W_2} \mathcal{F}(\mu), T_{\mu \rightarrow \mu_\star} - \text{id} \rangle_{\mu} - \frac{\alpha}{2} W_2^2(\mu, \mu_\star) \leq \frac{1}{2\alpha} \|\nabla_{W_2} \mathcal{F}(\mu)\|_{\mu}^2,$$

where the last line uses Young's inequality and $\|T_{\mu \rightarrow \mu_\star} - \text{id}\|_{\mu} = W_2(\mu, \mu_\star)$. However, (1.4.10) can hold even if $\mathcal{F}$ is not strongly convex.

- To convert convergence in function value to closeness to the minimizer, we note that (1.4.10) implies the following *quadratic growth* inequality for $\mathcal{F}$ around the minimizer:

$$\mathcal{F}(\nu) - \inf \mathcal{F} \geq \frac{\alpha}{2} W_2^2(\nu, \mu_\star) \quad \text{for all } \nu \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d). \quad (1.4.11)$$

It is easy to show that strong convexity implies (1.4.11); simply take $\mu = \mu_\star$ in (1.4.7), so that $\nabla_{W_2} \mathcal{F}(\mu_\star) = 0$. The proof of the stronger assertion—that (1.4.10) implies (1.4.11)—is sketched as [Exercise 1.16](#). In the case where $\mathcal{F}$ is not uniquely minimized, the right-hand side should be replaced by $\frac{\alpha}{2} \inf_{\mu_\star \in \arg \min \mathcal{F}} W_2^2(\nu, \mu_\star)$.

Finally, when $\mathcal{F}$ satisfies neither strong convexity nor the PL condition, (1.4.9) still implies convergence of the gradient norm:

$$\inf_{t \in [0, T]} \|\nabla_{W_2} \mathcal{F}(\mu_t)\|_{\mu_t} \leq \sqrt{\frac{\mathcal{F}(\mu_0) - \inf \mathcal{F}}{T}}.$$

In some cases, this can be used as the basis of a soft argument that $\mu_t \rightarrow \mu_\star$.

We now specialize these results to $\mathcal{F} = \text{KL}(\cdot \parallel \pi)$ with $\mu_\star = \pi \propto \exp(-V)$. Let $(\mu_t)_{t \geq 0}, (\nu_t)_{t \geq 0}$ denote the marginal laws for two copies of the Langevin diffusion (1.0.1). Then, α -convexity of V implies the bounds

$$W_2(\mu_t, \nu_t) \leq \exp(-\alpha t) W_2(\mu_0, \nu_0) \quad \text{and} \quad \text{KL}(\mu_t \parallel \pi) \leq \frac{\alpha}{2(\exp(\alpha t) - 1)} W_2^2(\mu_0, \pi).$$

In this setting, the first inequality is implied by the following theorem (see [Exercise 1.17](#)).

Theorem 1.4.12 (Contractivity of the Langevin diffusion). *Suppose that $\nabla^2 V \geq \alpha I_d$ for some $\alpha \in \mathbb{R}$. If $(X_t)_{t \geq 0}$ and $(X'_t)_{t \geq 0}$ denote two copies of the Langevin diffusion (1.0.1) with potential V and driven by the same Brownian motion, then we have the almost sure contraction*

$$\|X_t - X'_t\| \leq \exp(-\alpha t) \|X_0 - X'_0\|.$$

On the other hand, in this setting, the second inequality above is not sharp. The following is true:

$$\text{KL}(\mu_t \parallel \pi) \leq \frac{\alpha}{2(\exp(2\alpha t) - 1)} W_2^2(\mu_0, \pi). \quad (1.4.13)$$

A proof of this inequality is sketched in [Exercise 1.18](#) (see also [Exercise 2.17](#), [Exercise 3.4](#), and [Exercise 4.7](#)). In the case $\alpha = 0$, it yields the convergence rate

$$\text{KL}(\mu_t \parallel \pi) \leq \frac{1}{4t} W_2^2(\mu_0, \pi). \quad (1.4.14)$$

The PL inequality (1.4.10) specializes to

$$\text{KL}(\mu \parallel \pi) \leq \frac{1}{2\alpha} \left\| \nabla \log \frac{\mu}{\pi} \right\|_\mu^2 = \frac{1}{2\alpha} \text{FI}(\mu \parallel \pi) \quad \text{for all } \mu \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d). \quad (1.4.15)$$

This is precisely the log-Sobolev inequality (see [Example 1.2.27](#)); we have therefore recovered the Bakry–Émery theorem ([Theorem 1.2.30](#)) that $\text{CD}(\alpha, \infty)$ implies LSI, as well as [Theorem 1.2.26](#) which asserted that LSI yields exponentially fast decay in the KL divergence.

For the Langevin diffusion, the quadratic growth inequality (1.4.11) reads

$$\text{KL}(\mu \parallel \pi) \geq \frac{\alpha}{2} W_2^2(\mu, \pi) \quad \text{for all } \mu \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d). \quad (1.4.16)$$

This is known as **Talagrand's T_2 inequality** and it is an example of a *transport inequality*. Such inequalities have been closely studied in relation to the concentration of measure phenomenon. See [Chapter 2](#) and [van Handel \(2016, Chapter 4\)](#) for more details. The fact that LSI implies the T_2 inequality, which is a consequence of [Exercise 1.16](#), is known as the **Otto–Villani theorem** ([Otto and Villani, 2000](#)).

1.5 Overview of the Convergence Results

1.5.1 Convergence Results and Initialization

The main convergence results we have developed can be summarized as follows. Let $\alpha > 0$.

- $\text{KL}(\cdot \parallel \pi)$ is α -strongly convex along W_2 geodesics if and only if V is α -strongly convex, if and only if the following contraction holds. For all $\mu_0, \nu_0 \in \mathcal{P}_2(\mathbb{R}^d)$, if $(\mu_t)_{t \geq 0}, (\nu_t)_{t \geq 0}$ are Langevin diffusions started at μ_0 and ν_0 respectively, then $W_2^2(\mu_t, \nu_t) \leq \exp(-2\alpha t) W_2^2(\mu_0, \nu_0)$.

Moreover, in this case, $\text{KL}(\mu_t \parallel \pi) \leq \alpha W_2^2(\mu_0, \pi) / (2(\exp(2\alpha t) - 1))$, which then reduces to $\text{KL}(\mu_t \parallel \pi) \leq W_2^2(\mu_0, \pi) / (4t)$ for $\alpha = 0$.

- The target π satisfies the log-Sobolev inequality (LSI) with constant $1/\alpha$ if and only if for all $\mu_0 \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$, along the Langevin dynamics $(\mu_t)_{t \geq 0}$ started at μ_0 it holds that $\text{KL}(\mu_t \parallel \pi) \leq \exp(-2\alpha t) \text{KL}(\mu_0 \parallel \pi)$. The LSI is a gradient domination condition in Wasserstein space.
- The target π satisfies the Poincaré inequality (PI) with constant $1/\alpha$ if and only if for all $\mu_0 \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$, along the Langevin dynamics $(\mu_t)_{t \geq 0}$ started at μ_0 it holds that $\chi^2(\mu_t \parallel \pi) \leq \exp(-2\alpha t) \chi^2(\mu_0 \parallel \pi)$. The Poincaré inequality is a spectral gap condition for the generator of the Langevin diffusion.

The conditions are listed from strongest to weakest: α -strong log-concavity implies α^{-1} -LSI, which implies α^{-1} -PI. We also present convergence results which hold in Rényi divergences in Section 2.2.5.

In order to interpret these results, it is helpful to compare the sizes of these quantities to each other and at initialization.

Lemma 1.5.1 (Comparison between divergences). *For any $\mu \ll \pi$,*

$$2 \|\mu - \pi\|_{\text{TV}}^2 \leq \text{KL}(\mu \parallel \pi) \leq \log(1 + \chi^2(\mu \parallel \pi)) \leq \chi^2(\mu \parallel \pi),$$

where $\|\cdot\|_{\text{TV}}$ is the total variation distance, introduced in Definition 1.5.4.

Proof The first inequality is known as **Pinsker's inequality**. The second inequality is a special case of Exercise 2.9, and the last inequality follows from $\log(1 + x) \leq x$ for $x \geq 0$. $\square$

In the strongly log-concave case, the following initialization bound holds.

Lemma 1.5.2 (Chi-squared divergence at initialization). *Let $\pi \propto \exp(-V)$ be such that $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and $\kappa := \beta/\alpha$. Then, if $x_\star$ denotes the mode of π and $\mu_0 := \text{normal}(x_\star, \beta^{-1} I_d)$,*

$$\log(1 + \chi^2(\mu_0 \parallel \pi)) \leq \frac{d}{2} \log \kappa.$$

By Lemma 1.5.1, it implies that $\text{KL}(\mu_0 \parallel \pi) \leq \frac{d}{2} \log \kappa$. Recall that under a T_2 transport inequality with constant α^{-1} (which is implied by α^{-1} -LSI) we have $\text{KL} \geq \frac{\alpha}{2} W_2^2$, so this further implies $W_2^2(\mu_0, \pi) \leq (d/\alpha) \log \kappa$. This inequality can be sharpened slightly.

Lemma 1.5.3 (Wasserstein distance at initialization). *Let $\pi \propto \exp(-V)$ be such that $0 < \alpha I_d \leq \nabla^2 V$. Let $x_\star$ denote the mode of π . Then, for any μ_0 with mean m and covariance Σ ,*

$$W_2^2(\mu_0, \pi) \leq \frac{d}{\alpha} + \text{tr } \Sigma + \|m - x_\star\| \left(\|m - x_\star\| + 2\sqrt{\frac{d}{\alpha}} \right).$$

Proof Let $X_0 \sim \mu_0$ and $Z \sim \pi$ be independent. Then, by conditioning on Z ,

$$\begin{aligned} W_2^2(\mu_0, \pi) &\leq \mathbb{E}[\|X_0 - Z\|^2] = \text{tr } \Sigma + \mathbb{E}[\|Z - m\|^2] \\ &= \text{tr } \Sigma + \mathbb{E}[\|m - x_\star\|^2 + \|Z - x_\star\|^2 - 2\langle m - x_\star, Z - x_\star \rangle] \\ &\leq \text{tr } \Sigma + \mathbb{E}[\|Z - x_\star\|^2] + \|m - x_\star\| (\|m - x_\star\| + 2\sqrt{\mathbb{E}[\|Z - x_\star\|^2]}) \\ &\leq \text{tr } \Sigma + \frac{d}{\alpha} + \|m - x_\star\| (\|m - x_\star\| + 2\sqrt{\frac{d}{\alpha}}), \end{aligned}$$

where the bound on $\mathbb{E}[\|Z - x_\star\|^2]$ follows from [Lemma 4.0.1](#). $\square$

In particular, for $\mu_0 = \delta_{x_\star}$, it yields $W_2^2(\delta_{x_\star}, \pi) \leq d/\alpha$.

Therefore, at initialization, we typically have $W_2^2(\mu_0, \pi)$, $\text{KL}(\mu_0 \parallel \pi) = \tilde{O}(d)$ (at least when π is strongly log-concave). Hence, exponential convergence in W_2^2 and KL both imply that the amount of time it takes for the Langevin diffusion to reach ε error is $O(\log(d/\varepsilon))$. In contrast, the chi-squared divergence is typically much larger at initialization: $\chi^2(\mu_0 \parallel \pi) = \exp(O(d))$. Therefore, the chi-squared result implies that the Langevin diffusion takes $O(d \vee \log(1/\varepsilon))$ time to reach ε error.

When we study discretization in Chapter 4, the W_2 contraction under strong log-concavity is the easiest to turn into a sampling guarantee. This is because to bound the W_2 distance, we can use straightforward *coupling* techniques. On the other hand, a continuous-time result in KL or χ^2 often requires the discretization analysis to also be carried out in KL or χ^2 , which is substantially trickier.

For more general initialization results, see [Chewi et al. \(2025a, Appendix A\)](#).

1.5.2 Appendix: Divergences between Probability Measures

As we have already seen, the analysis of Langevin introduces many different notions of divergences between probability measures. Therefore, it is important to develop a healthy understanding of the relationships between these divergences.

First of all, there is a distinction between the Wasserstein metric, which is a transport distance (measuring how far we must move the mass of one measure to the other), and information divergences which are defined directly in terms of the densities such as the KL divergence and the chi-squared divergence. Note that the latter two divergences are infinite unless the first argument is absolutely continuous w.r.t. the second, which is certainly not the case for the Wasserstein metric.

We introduce another important metric.

Definition 1.5.4. The **total variation (TV) distance** between probability measures $\mu, \nu \in \mathcal{P}(\mathcal{X})$ is defined via

$$\begin{aligned} \|\mu - \nu\|_{\text{TV}} &:= \sup_{A \subseteq \mathcal{X}} |\mu(A) - \nu(A)| = \sup_{f: \mathcal{X} \rightarrow [0,1]} \left| \int f \, d\mu - \int f \, d\nu \right| \\ &= \inf_{\gamma \in \mathcal{C}(\mu, \nu)} \int \mathbb{1}\{x \neq y\} \gamma(\text{d}x, \text{d}y) = \frac{1}{2} \int \left| \frac{\text{d}\mu}{\text{d}\lambda} - \frac{\text{d}\nu}{\text{d}\lambda} \right| \text{d}\lambda, \end{aligned}$$

where λ is a common dominating measure for μ and ν .

The TV metric is indeed a metric on the space $\mathcal{P}(\mathcal{X})$ (in fact, it can be extended to a norm on the space $\mathcal{M}(\mathcal{X})$ of signed measures). The TV distance can be thought of as both a transport metric

(with cost $(x, y) \mapsto \mathbb{1}\{x \neq y\}$; in fact, the TV distance is a special case of the W_1 metric introduced in [Exercise 1.10](#)) and an information divergence.

The family of information divergences can be further expanded as follows ([Csiszár, 1967](#)).

Definition 1.5.5. Let $\mu, \nu \in \mathcal{P}(\mathcal{X})$, and let $f : \mathbb{R}_+ \rightarrow \mathbb{R}$ be a convex function with $f(1) = 0$. Then, the f -divergence of μ relative to ν is

$$\mathcal{D}_f(\mu \parallel \nu) := \int f\left(\frac{d\mu}{d\nu}\right) d\nu, \quad \text{if } \mu \ll \nu.$$

In general, if $\mu \not\ll \nu$, we let p_μ, p_ν denote the respective densities of μ and ν w.r.t. a common dominating measure. Then, for $f'(\infty) := \lim_{t \rightarrow \infty} f(t)/t$,

$$\mathcal{D}_f(\mu \parallel \nu) := \int_{p_\nu > 0} f\left(\frac{p_\mu}{p_\nu}\right) d\nu + f'(\infty) \mu\{p_\nu = 0\}.$$

For example, the TV distance corresponds to $f(x) = \frac{1}{2}|x - 1|$, the KL divergence corresponds to $f(x) = x \log x$, and the χ^2 divergence corresponds to $f(x) = (x - 1)^2$. When f has superlinear growth, then $f'(\infty) = \infty$ and hence $\mathcal{D}_f(\mu \parallel \nu) = \infty$ unless $\mu \ll \nu$, but the second more general definition given above is necessary to recover the TV distance.

In [Section 2.2.5](#), we will introduce the closely related family of Rényi divergences.

We conclude by stating a few key facts (without complete proofs) about f -divergences. The first is the data-processing inequality. Below, given a Markov kernel P and a mixing measure μ , we write $\mu P := \int P(x, \cdot) \mu(dx)$ for the mixture distribution.

Theorem 1.5.6 (Data-processing inequality). *Suppose that $\mu, \nu \in \mathcal{P}(\mathcal{X})$ and that P is any Markov kernel. Then, for any f -divergence, it holds that*

$$\mathcal{D}_f(\mu P \parallel \nu P) \leq \mathcal{D}_f(\mu \parallel \nu).$$

Consequently, $\mathcal{D}_f(\cdot \parallel \cdot)$ is jointly convex in its two arguments.

Proof For $X \sim \nu, Y \sim P(X, \cdot)$, we have the identity

$$\frac{d(\mu P)}{d(\nu P)}(Y) = \mathbb{E}\left[\frac{d\mu}{d\nu}(X) \mid Y\right].$$

The data-processing inequality now immediately follows from Jensen's inequality, by applying f and taking expectations.

To prove joint convexity, let $\lambda \in (0, 1)$ and $Z \sim \text{Bernoulli}(\lambda)$. Let $\mu_0, \mu_1, \nu_0, \nu_1 \in \mathcal{P}(\mathcal{X})$. Under μ , let X be drawn from μ_Z , and under ν , let X be drawn from ν_Z . Then,

$$\mathcal{D}_f(\text{law}_\mu(X, Z) \parallel \text{law}_\nu(X, Z)) = (1 - \lambda) \mathcal{D}_f(\mu_0 \parallel \nu_0) + \lambda \mathcal{D}_f(\mu_1 \parallel \nu_1).$$

Also, $\mathcal{D}_f(\text{law}_\mu(X, Z) \parallel \text{law}_\nu(X, Z)) \geq \mathcal{D}_f(\text{law}_\mu(X) \parallel \text{law}_\nu(X))$ by the data-processing inequality, where we note that

$$\mathcal{D}_f(\text{law}_\mu(X) \parallel \text{law}_\nu(X)) = \mathcal{D}_f((1 - \lambda) \mu_0 + \lambda \mu_1 \parallel (1 - \lambda) \nu_0 + \lambda \nu_1). \quad \square$$

The remaining facts are specific to the KL divergence. The Donsker–Varadhan theorem expresses the KL divergence via a variational principle.

Theorem 1.5.7 (Donsker–Varadhan variational principle). *Suppose that $\mu, \nu \in \mathcal{P}(\mathcal{X})$, where $\mathcal{X}$ is a Polish space. Then,*

$$\mathrm{KL}(\mu \parallel \nu) = \sup \{ \mathbb{E}_\mu g - \log \mathbb{E}_\nu \exp g \mid g : \mathcal{X} \rightarrow \mathbb{R} \text{ is bounded and measurable} \}.$$

The theorem asserts that the functionals $\mu \mapsto \mathrm{KL}(\mu \parallel \nu)$ and $g \mapsto \log \mathbb{E}_\nu \exp g$ are convex conjugates of each other. See [Dembo and Zeitouni \(2010, Lemma 6.2.13\)](#) or [Rassoul-Agha and Seppäläinen \(2015, Theorem 5.4\)](#) for careful proofs, or see the remark after [Lemma 2.3.5](#).

Lastly, we have the chain rule for the KL divergence.

Lemma 1.5.8 (Chain rule for KL divergence). *Let $\mathcal{X}_1, \mathcal{X}_2$ be Polish spaces and suppose we are given two probability measures $\mu, \nu \in \mathcal{P}(\mathcal{X}_1 \times \mathcal{X}_2)$ with $\mu \ll \nu$. Let μ_1 be the $\mathcal{X}_1$ marginal of μ , and let $\mu_{2|1}(\cdot \mid \cdot)$ be the conditional distribution for μ on $\mathcal{X}_2$ conditioned on $\mathcal{X}_1$; likewise define ν_1 and $\nu_{2|1}$. Then, it holds that*

$$\mathrm{KL}(\mu \parallel \nu) = \mathrm{KL}(\mu_1 \parallel \nu_1) + \int \mathrm{KL}(\mu_{2|1}(\cdot \mid x_1) \parallel \nu_{2|1}(\cdot \mid x_1)) \mu_1(dx_1).$$

We invite the reader to prove the chain rule in the discrete case ($\mathcal{X}_1$ and $\mathcal{X}_2$ are finite sets), free of measure-theoretic guilt.

Bibliographical Notes

Much of the material in this chapter is foundational, with entire textbooks giving comprehensive treatments of the topics. For stochastic calculus, there is of course a long list of textbooks, but as a starting place we suggest [Steele \(2001\)](#); [Pavliotis \(2014\)](#); [Le Gall \(2016\)](#). For Markov semigroup theory, see [Bakry et al. \(2014\)](#); [van Handel \(2016\)](#). For optimal transport, the core theory is developed in [Villani \(2003\)](#); [Santambrogio \(2015\)](#); [Chewi et al. \(2025b\)](#), and for a rigorous development of Otto calculus see [Ambrosio et al. \(2008\)](#); [Villani \(2009b\)](#).

The notion of solution used in Section 1.1.3 is more typically called a *strong solution* to the SDE, because given any Brownian motion B we can find a process X which is driven by B and which satisfies the SDE. There is also a notion of *weak solution*, in which we are allowed to construct the probability space $(\Omega, \mathcal{F}, (\mathcal{F}_t)_{t \geq 0}, \mathbb{P})$ and the Brownian motion B together with the solution X . We will not worry about the distinction in this book, since strong solutions suffice for our purposes.

In accordance with most of the literature, we refer to the last conclusion of [Theorem 1.3.8](#) as Brenier’s theorem, after [Brenier \(1991\)](#). However, see also the earlier work of [Cuesta and Matrán \(1989\)](#). An elegant direct proof of strong duality can be found in [Santambrogio \(2015, Section 1.6.3\)](#).

The perspective of the Langevin diffusion as a Wasserstein gradient flow was introduced in [Jordan et al. \(1998\)](#); the application of Otto calculus to functional inequalities was given in [Otto and Villani \(2000\)](#); and the calculation rules for Otto calculus were set out in [Otto \(2001\)](#). These three papers are seminal and are worth reading carefully. An alternative (but related) approach to functional inequalities via optimal transport is given in [Cordero-Erausquin \(2002\)](#). The use of PL inequalities in optimization was pioneered in [Polyak \(1963\)](#) and further popularized by [Karimi et al. \(2016\)](#). The formal proof of the Otto–Villani theorem in [Exercise 1.16](#) was made rigorous via entropic interpolations in [Gentil et al. \(2020\)](#); see [Blanchet and Bolte \(2018\)](#) for a generalization.

The Efron–Stein inequality in [Exercise 1.1](#), due to [Efron and Stein \(1981\)](#), is just one example of

the use of martingales to derive concentration inequalities; see [Boucheron et al. \(2013\)](#); [van Handel \(2016\)](#) for more on this topic.

The upper bound (1.E.2) in [Exercise 1.9](#) is surprisingly sharp: it holds that

$$\frac{1}{2} \|\Sigma_0^{1/2} - \Sigma_1^{1/2}\|_{\text{HS}}^2 \leq W_2^2(\mu_0, \mu_1) - \|m_0 - m_1\|^2 \leq \|\Sigma_0^{1/2} - \Sigma_1^{1/2}\|_{\text{HS}}^2,$$

see [Carrillo and Vaes \(2021, Lemma 3.5\)](#).

The proof of the dynamical formulation of dual optimal transport in [Exercise 1.11](#) is carried out rigorously in [Villani \(2003, Section 8.1\)](#). We mention that the Hamilton–Jacobi equation has a close connection with classical mechanics; in particular, the characteristics of the Hamilton–Jacobi equation are precisely Hamilton’s equations of motion ([Evans, 2010, Section 3.3](#)). In the context of optimal transport, the Hamiltonian consists only of kinetic energy (no potential energy) and hence the characteristics are straight lines traversed at constant speed; this is consistent with the description of Wasserstein geodesics. The Hamilton–Jacobi equation, the Hopf–Lax semigroup, and their connection with optimal transport can also be generalized to other costs; see [Villani \(2003, Section 5.4\)](#).

The proof of the sharp KL convergence bound in [Exercise 1.18](#) was carried out for $\alpha = 0$ in [Otto and Villani \(2001\)](#), and for $\alpha > 0$ in [Liang et al. \(2024\)](#).

Exercises

A Primer on Stochastic Calculus

▷ [Exercise 1.1](#) (Orthogonality of martingale increments)

Let $(M_n)_{n \in \mathbb{N}}$ be a discrete-time martingale which is adapted to a filtration $(\mathcal{F}_n)_{n \in \mathbb{N}}$ and satisfies $\mathbb{E}[M_n^2] < \infty$ for all $n \in \mathbb{N}$. Let $\Delta_n := M_{n+1} - M_n$ denote the martingale increment.

- 1 Prove that for $m, n \in \mathbb{N}$ with $m \neq n$, $\mathbb{E}[\Delta_m \Delta_n] = 0$: the martingale increments are orthogonal. In particular, if $M_0 = 0$, then $\mathbb{E}[M_n^2] = \sum_{k=0}^{n-1} \mathbb{E}[\Delta_k^2]$.
- 2 Let $(X_i)_{i=1}^n$ be independent random variables taking values in some space $\mathcal{X}$, and suppose that the function $f : \mathcal{X}^n \rightarrow \mathbb{R}$ is bounded and measurable. Check that if $M_k := \mathbb{E}[f(X_1, \dots, X_n) \mid X_1, \dots, X_k]$, then the **Doob martingale** $(M_k)_{k=0}^n$ is indeed a martingale. Then, using the previous part, prove the following tensorization property of the variance:

$$\text{var } f(X_1, \dots, X_n) \leq \mathbb{E} \sum_{k=1}^n \text{var}(f(X_1, \dots, X_n) \mid X_{-k}),$$

where $X_{-k} := (X_1, \dots, X_{k-1}, X_{k+1}, \dots, X_n)$.

- 3 Define the discrete derivative

$$D_k f(x) := \sup_{x'_k \in \mathcal{X}} f(x_1, \dots, x_{k-1}, x'_k, x_{k+1}, \dots, x_n) - \inf_{x'_k \in \mathcal{X}} f(x_1, \dots, x_{k-1}, x'_k, x_{k+1}, \dots, x_n).$$

Prove the inequality

$$\text{var } f(X_1, \dots, X_n) \leq \frac{1}{4} \mathbb{E} \sum_{k=1}^n \{D_k f(X_1, \dots, X_n)\}^2.$$

This inequality, known as the **Efron–Stein inequality**, expresses the fact that a function f of independent random variables which is not too sensitive to any individual coordinate has controlled

variance. This is a concentration inequality which has useful consequences in many probabilistic settings, see, e.g., [Boucheron et al. \(2013\)](#).

Hint: First prove that a random variable which takes values in $[a, b]$ has variance bounded by $\frac{1}{4}(b-a)^2$.

▷ **Exercise 1.2 (Explosion of ODEs)**

Solve the following ODE on $\mathbb{R}$: $\dot{x}_t = b(x_t)$ with initial condition $x_0 \in \mathbb{R}$, where $b(x) = |x|^\alpha$, $\alpha > 0$. Show that

- 1 when $0 < \alpha < 1$, there are multiple solutions to the ODE (with initial condition $x_0 = 0$) so that uniqueness fails;
- 2 when $\alpha = 1$ (and hence b is globally Lipschitz), there is a unique solution to the ODE which is finite for all time;
- 3 when $\alpha > 1$, then the solution to the ODE blows up in finite time, provided $x_0 > 0$.

▷ **Exercise 1.3 (Ornstein–Uhlenbeck process)**

One of the most important diffusions that we will encounter is the **Ornstein–Uhlenbeck (OU) process**, which solves the SDE

$$dX_t = -X_t dt + \sqrt{2} dB_t.$$

Give an explicit expression for X_t in terms of X_0 and an Itô integral involving $(B_t)_{t \geq 0}$. From this expression, can you read off the stationary distribution of this process?

Hint: Apply Itô's formula to $f(t, X_t) = X_t \exp t$. To find the stationary distribution, justify the following fact: if $(\eta_t)_{t \geq 0}$ is a *deterministic* function, then $\int_0^T \eta_t dB_t$ is a Gaussian with mean zero and variance $\int_0^T \eta_t^2 dt$.

Markov Semigroup Theory

▷ **Exercise 1.4 (Generator for general SDEs)**

Compute the generator for a general time-homogeneous SDE $dX_t = b(X_t) dt + \sigma(X_t) dB_t$.

▷ **Exercise 1.5 (Basic properties of the Markov semigroup)**

Let $(P_t)_{t \geq 0}$ be a Markov semigroup with carré du champ Γ .

- 1 Prove that if $\phi : \mathbb{R} \rightarrow \mathbb{R}$ is convex, then $P_t \phi(f) \geq \phi(P_t f)$ and $\mathcal{L} \phi(f) \geq \phi'(f) \mathcal{L} f$ whenever the expressions are well-defined.
- 2 If $(X_t)_{t \geq 0}$ denotes the Markov process associated with the semigroup, assumed to have continuous sample paths, and f is smooth, then the process $t \mapsto f(X_t) - \int_0^t \mathcal{L} f(X_s) ds$ is a continuous local martingale. In particular, $(f(X_t))_{t \geq 0}$ is a continuous local martingale for all initializations if and only if $\mathcal{L} f = 0$.
- 3 Prove that for $f, g \in L^2(\pi)$ in the domain of the carré du champ, we have the Cauchy–Schwarz inequality $|\Gamma(f, g)| \leq \sqrt{\Gamma(f, f) \Gamma(g, g)}$.

Hint: For $\lambda \in \mathbb{R}$, consider $0 \leq \Gamma(f + \lambda g, f + \lambda g)$ and use bilinearity.

▷ **Exercise 1.6 (Functional inequalities and exponential decay)**

Prove the equivalence between the Poincaré inequality and exponential decay of variance ([Theorem 1.2.21](#)), and the equivalence between the log-Sobolev inequality and exponential decay of entropy ([Theorem 1.2.26](#)).

▷ **Exercise 1.7** (Log-Sobolev implies Poincaré)

Linearize the log-Sobolev inequality to obtain the Poincaré inequality.

Hint: Argue that if $f \in C_c^\infty(\mathbb{R}^d)$ satisfies $\int f \, d\pi = 0$, then

$$\text{KL}((1 + \varepsilon f) \pi \parallel \pi) = \frac{\varepsilon^2}{2} \int f^2 \, d\pi + o(\varepsilon^2). \quad (1.E.1)$$

▷ **Exercise 1.8** (Mixing of the Ornstein–Uhlenbeck process)

Consider the Ornstein–Uhlenbeck process $(X_t)_{t \geq 0}$ introduced in [Exercise 1.3](#). Note that this is just an instance of the Langevin diffusion with potential $V(x) = \frac{\|x\|^2}{2}$.

1 Using the explicit solution of the OU process, show that the semigroup has the explicit expression

$$P_t f(x) = \mathbb{E} f(\exp(-t)x + \sqrt{1 - \exp(-2t)} \xi), \quad \xi \sim \text{normal}(0, I_d).$$

Using this expression for the semigroup, compute the generator by hand and check that it agrees with the general formula obtained in [Example 1.2.4](#).

2 Show that for the OU process, $\nabla P_t f = \exp(-t) P_t \nabla f$. Next, show that

$$\mathcal{E}(P_t f, P_t f) \leq \exp(-2t) \mathcal{E}(f, f).$$

Explain why this implies a Poincaré inequality for the standard Gaussian distribution.

Hint: For a general Markov semigroup, show that $\text{var}_\pi f = 2 \int_0^\infty \mathcal{E}(P_t f, P_t f) \, dt$ by differentiating $t \mapsto \text{var}_\pi(P_t f)$.

In Chapter 2, we will generalize these calculations to prove [Theorem 1.2.30](#).

The Geometry of Optimal Transport

▷ **Exercise 1.9** (Optimal transport between Gaussians)

Let $\mu_0 := \text{normal}(m_0, \Sigma_0)$ and $\mu_1 := \text{normal}(m_1, \Sigma_1)$; assume that $\Sigma_0 > 0$. Compute the optimal transport map from μ_0 to μ_1 , as well as the cost $W_2(\mu_0, \mu_1)$. [By Brenier’s theorem, it suffices to find the gradient of a convex function which pushes forward μ_0 to μ_1 .]

Also, exhibit a coupling to prove the upper bound

$$W_2^2(\mu_0, \mu_1) \leq \|m_0 - m_1\|^2 + \|\Sigma_0^{1/2} - \Sigma_1^{1/2}\|_{\text{HS}}^2. \quad (1.E.2)$$

Finally, suppose that ν_0, ν_1 are probability measures, and suppose that μ_0, μ_1 are Gaussians whose means and covariances match those of ν_0 and ν_1 respectively. Then, prove that $W_2(\nu_0, \nu_1) \geq W_2(\mu_0, \mu_1)$.

Hint: For the last statement, use the fact that the dual potentials for optimal transport between Gaussians are quadratic functions.

▷ **Exercise 1.10** (Optimal transport with other costs)

In this exercise, we consider optimal transport with a general cost function c as in (1.3.2).

1 By following the proof of [Theorem 1.3.8](#), argue that if optimal dual potentials (f, g) exist, they can be taken to be c -conjugates, i.e., $f = g^c$ and $g = f^c$ where

$$g^c(x) := \inf_{y \in \mathcal{Y}} \{c(x, y) - g(y)\}, \quad f^c(y) := \inf_{x \in \mathcal{X}} \{c(x, y) - f(x)\}, \quad (1.E.3)$$

and that

$$\mathcal{T}_c(\mu, \nu) \geq \sup_{f \in L^1(\mu)} \left\{ \int f \, d\mu + \int f^c \, d\nu \right\}.$$

Under general conditions, equality holds; see Villani (2009b, Theorem 5.10).

Functions of the form (1.E.3) are called **c-concave**.

2 Let $\mathcal{X} = \mathcal{Y}$ be a metric space with metric d . For all $p \geq 1$, we can define the **p -Wasserstein distance**

$$W_p^p(\mu, \nu) = \inf_{\gamma \in \mathcal{C}(\mu, \nu)} \int d(x, y)^p \gamma(dx, dy).$$

Let $\mathcal{P}_p(\mathcal{X})$ denote the space of probability measures μ over $\mathcal{X}$ such that for some $x_0 \in \mathcal{X}$, $\int d(x_0, \cdot)^p \, d\mu < \infty$. Show that $(\mathcal{P}_p(\mathcal{X}), W_p)$ is a metric space.

3 In the case $p = 1$, show that if (f, g) are d -conjugates, then $f = -g$ and f is 1-Lipschitz. Deduce the duality formula

$$W_1(\mu, \nu) = \sup \left\{ \int f \, d\mu - \int f \, d\nu \mid f : \mathcal{X} \rightarrow \mathbb{R} \text{ is 1-Lipschitz} \right\}. \quad (1.E.4)$$

▷ **Exercise 1.11 (Dynamical formulations of optimal transport)**

The formula (1.3.24) shows that the W_2 distance between μ_0 and μ_1 equals the smallest arc length of any curve joining μ_0 and μ_1 . It is also true that the squared W_2 distance minimizes the *energy* or *action* of any curve joining μ_0 and μ_1 , in the following sense:

$$W_2^2(\mu_0, \mu_1) = \inf \left\{ \int_0^1 \|v_t\|_{L^2(\mu_t)}^2 \, dt \mid \partial_t \mu_t + \operatorname{div}(\mu_t v_t) = 0 \right\}. \quad (1.E.5)$$

Although the problems (1.3.24) and (1.E.5) both identify geodesics in the Wasserstein space, the latter problem has more favorable properties. Namely, the minimizing curves in (1.3.24) are geodesics, but they may not have constant speed (indeed, the arc length functional is invariant under time reparameterization of the curve); in contrast, minimizing curves in (1.E.5) necessarily have constant speed. Also, we can reparameterize problem (1.E.5) by introducing the *momentum density* $p_t := \mu_t v_t$ and write

$$W_2^2(\mu_0, \mu_1) = \inf \left\{ \int_0^1 \left(\int \frac{\|p_t(x)\|^2}{\mu_t(x)} \, dx \right) \, dt \mid \partial_t \mu_t + \operatorname{div} p_t = 0 \right\}, \quad (1.E.6)$$

which is now a convex problem in the variables (μ, p) . This convenient reformulation is known as the **Benamou–Brenier formula** (Benamou and Brenier, 2000).

Just as (1.E.5) describes the *dynamical* version of the *static* optimal transport problem (1.3.5), there is a dynamical formulation of the dual optimal transport problem (1.3.7), in which the dual potential evolves according to the **Hamilton–Jacobi equation**

$$\partial_t u_t + \frac{1}{2} \|\nabla u_t\|^2 = 0. \quad (1.E.7)$$

Then, it holds that

$$\frac{1}{2} W_2^2(\mu_0, \mu_1) = \sup \left\{ \int u_1 \, d\mu_1 - \int u_0 \, d\mu_0 \mid \partial_t u_t + \frac{1}{2} \|\nabla u_t\|^2 = 0 \right\}. \quad (1.E.8)$$

The goal of this exercise is to justify and understand these facts.

- 1 Show that the mapping $\mathbb{R}_{>0} \times \mathbb{R}^d \rightarrow \mathbb{R}$, $(\mu, p) \mapsto \|p\|^2/\mu$ is convex. Also, compute the convex conjugate of this mapping. Deduce that the Benamou–Brenier reformulation (1.E.6) is convex.
- 2 Ignoring issues of regularity, show that the solution u_t of the Hamilton–Jacobi equation with initial condition $u_0 = f$ is described by the **Hopf–Lax semigroup**

$$u_t(x) = Q_t f(x) := \inf_{y \in \mathbb{R}^d} \left\{ f(y) + \frac{1}{2t} \|y - x\|^2 \right\}.$$

- 3 Following the proof of Theorem 1.3.8, show that the dual optimal transport problem (1.3.7) can be written

$$\frac{1}{2} W_2^2(\mu_0, \mu_1) = \sup_{f \in L^1(\mu_0)} \left\{ \int Q_1 f \, d\mu_1 - \int f \, d\mu_0 \right\},$$

where Q_1 denotes the Hopf–Lax semigroup at time 1. From this, deduce that the formula (1.E.8) holds.

- 4 Although the previous part gives a proof of the dynamical formulation (1.E.8), it is unsatisfactory because it only involves an analysis of the static primal and dual problems. Here, we present a purely dynamical proof. The continuity constraint $\partial_t \mu_t + \operatorname{div} p_t = 0$ in (1.E.6) can be reformulated as follows: for any curve of functions $[0, 1] \times \mathbb{R}^d \rightarrow \mathbb{R}$, $(t, x) \mapsto u_t(x)$,

$$\begin{aligned} \int u_1 \, d\mu_1 - \int u_0 \, d\mu_0 &= \int_0^1 \left(\partial_t \int u_t \, d\mu_t \right) dt = \int_0^1 \left(\int (\partial_t u_t \mu_t + u_t \partial_t \mu_t) \right) dt \\ &= \int_0^1 \left(\int (\partial_t u_t + \langle \nabla u_t, \frac{p_t}{\mu_t} \rangle) \, d\mu_t \right) dt. \end{aligned}$$

This can be incorporated as a Lagrange multiplier in (1.E.6):

$$\begin{aligned} \frac{1}{2} W_2^2(\mu_0, \mu_1) &= \inf_{\substack{\mu: [0,1] \times \mathbb{R}^d \rightarrow \mathbb{R}_+ \\ p: [0,1] \times \mathbb{R}^d \rightarrow \mathbb{R}^d}} \sup_{u: [0,1] \times \mathbb{R}^d \rightarrow \mathbb{R}} \left\{ \int_0^1 \int \frac{\|p_t\|^2}{2\mu_t} \, dt + \int u_1 \, d\mu_1 - \int u_0 \, d\mu_0 \right. \\ &\quad \left. - \int_0^1 \left(\int (\partial_t u_t + \langle \nabla u_t, \frac{p_t}{\mu_t} \rangle) \, d\mu_t \right) dt \right\} \end{aligned}$$

Assume that the infimum and supremum can be interchanged (here, we are invoking an abstract minimax theorem, which crucially relies on the convexity of the problem established in the first part). Use this to prove that

$$\frac{1}{2} W_2^2(\mu_0, \mu_1) = \sup \left\{ \int u_1 \, d\mu_1 - \int u_0 \, d\mu_0 \mid \partial_t u_t + \frac{1}{2} \|\nabla u_t\|^2 \leq 0 \right\}$$

and that equality holds only if the Hamilton–Jacobi equation (1.E.7) holds, and if $\nabla u_t = v_t = p_t/\mu_t$. Note that this also establishes that the optimal vector fields $(v_t)_{t \in [0,1]}$ are gradients of functions.

- 5 Let $(\mu_t)_{t \in [0,1]}$ be a Wasserstein geodesic and let $(v_t)_{t \in [0,1]}$ be its associated curve of tangent vectors. Prove that $\partial_t v_t + \nabla_{v_t} v_t = 0$. (This statement follows from the previous part by differentiating the Hamilton–Jacobi equation in space. Try to also give a more direct proof of this equation.)

Hint: If $\dot{x}_t = v_t(x_t)$, then because particles travel with constant velocity along Wasserstein geodesics, $t \mapsto \dot{x}_t$ is constant.

⊇ **Exercise 1.12** (Wasserstein space has non-negative curvature)

Let $(\mu_t)_{t \in [0,1]}$ denote a Wasserstein geodesic. By finding an appropriate coupling, prove that for all $t \in [0, 1]$ and all $\nu \in \mathcal{P}_2(\mathbb{R}^d)$,

$$W_2^2(\mu_t, \nu) \geq (1-t) W_2^2(\mu_0, \nu) + t W_2^2(\mu_1, \nu) - t(1-t) W_2^2(\mu_0, \mu_1). \quad (1.E.9)$$

Compare this to the following equality on $\mathbb{R}^d$: if $x_t = (1-t)x_0 + tx_1$, then

$$\|x_t - y\|^2 = (1-t) \|x_0 - y\|^2 + t \|x_1 - y\|^2 - t(1-t) \|x_0 - x_1\|^2. \quad (1.E.10)$$

The equality (1.E.10) expresses the fact that $\mathbb{R}^d$ is *flat*, whereas the inequality (1.E.9) expresses the fact that $\mathcal{P}_2(\mathbb{R}^d)$ (equipped with the W_2 metric) is *non-negatively curved*, like a sphere; see Ambrosio et al. (2008, Section 7.3).

The Langevin SDE as a Wasserstein Gradient Flow

▷ Exercise 1.13 (Reconciling the SDE and Wasserstein perspectives)

Let $\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$ be a smooth test function and let $\delta > 0$. First, consider the Langevin diffusion $dZ_t = -\nabla V(Z_t) dt + \sqrt{2} dB_t$ started at $Z_0 \sim \mu_0$, and compute $\mathbb{E} \varphi(Z_\delta)$ up to first order in δ .

Next, if $\partial_t \mu_t = \operatorname{div}(\mu_t \nabla \log(\mu_t/\pi))$, with $\pi \propto \exp(-V)$, then we can interpret this as a fluid flow: let $X_0 \sim \mu_0$, and $\dot{X}_t = -\nabla \log(\mu_t/\pi)(X_t)$, so that $X_t \sim \mu_t$. Compute $\mathbb{E} \varphi(X_\delta)$ up to first order in δ .

Check that the two expressions you computed match (up to first order in δ). Note that these calculations are implicit in Sections 1.2 and 1.4, but it is illuminating to directly connect the Langevin diffusion to the Wasserstein gradient flow.

▷ Exercise 1.14 (Wasserstein calculus for f -divergences)

Compute the Wasserstein gradient of the functional $\chi^2(\cdot \| \pi)$. Use the rules of Wasserstein calculus to compute $\partial_t \chi^2(\pi_t \| \pi)$, where $(\pi_t)_{t \geq 0}$ is the law of the Langevin diffusion with stationary distribution π . Check that the result agrees with a calculation based on Markov semigroup theory.

More generally, let $f : \mathbb{R}_+ \rightarrow \mathbb{R}$ be convex with $f(1) = 0$, and consider the corresponding f -divergence $\mathcal{D}_f$ (Definition 1.5.5). Compute the Wasserstein gradient of $\mathcal{D}_f(\cdot \| \pi)$. For bonus points, calculate the Wasserstein Hessian as well.

▷ Exercise 1.15 (Smoothness along Wasserstein geodesics)

For a functional $\mathcal{F}$ over the Wasserstein space, let us say that it is β -smooth if, for all constant-speed geodesics $(\mu_t)_{t \in [0,1]}$ with initial tangent vector $v_0 = T - \operatorname{id}$, it holds that $\partial_t^2|_{t=0} \mathcal{F}(\mu_t) \leq \beta \|T - \operatorname{id}\|_{\mu_0}^2$.

- 1 Show that the potential energy functional $\mathcal{E}$ corresponding to a potential V with $\nabla^2 V \leq \beta I_d$ is β -smooth.
- 2 Establish the expression (1.4.4) for the entropy $\mathcal{H}$ and argue that $\mathcal{H}$ is non-smooth.

▷ Exercise 1.16 (Otto–Villani theorem)

Consider the gradient flow $(\mu_t)_{t \geq 0}$ of a functional $\mathcal{F}$ with $\inf \mathcal{F} = 0$. Assume that the PL inequality $\|\nabla_{W_2} \mathcal{F}(\mu)\|_\mu^2 \geq 2\alpha \mathcal{F}(\mu)$ holds with $\alpha > 0$, and that the gradient flow converges to the minimizer of $\mathcal{F}$. Argue that $\partial_t W_2(\mu_t, \mu_0) \leq \|\nabla_{W_2} \mathcal{F}(\mu_t)\|_{\mu_t}$, and then show that

$$\partial_t \left(\sqrt{\frac{\alpha}{2}} W_2(\mu_t, \mu_0) + \sqrt{\mathcal{F}(\mu_t)} \right) \leq 0.$$

Conclude that a quadratic growth inequality holds.

▷ **Exercise 1.17** (Contraction of the Langevin diffusion)

In this exercise, we explore different proofs of contraction.

- 1 Suppose that $V : \mathbb{R}^d \rightarrow \mathbb{R}$ is α -strongly convex. Let $(x_t)_{t \geq 0}, (y_t)_{t \geq 0}$ be two gradient flows for V . Show that $\|x_t - y_t\|^2 \leq \exp(-2\alpha t) \|x_0 - y_0\|^2$.
- 2 Next, prove [Theorem 1.4.12](#). *Hint:* Apply Itô's formula ([Theorem 1.1.19](#)) to $f(x, x') := \|x - x'\|^2$.
- 3 Let $\mathcal{F} : \mathcal{P}_{2,ac}(\mathbb{R}^d) \rightarrow \mathbb{R}$ be an α -convex functional, and let $(\mu_t)_{t \geq 0}, (\nu_t)_{t \geq 0}$ be gradient flows for $\mathcal{F}$. Prove that

$$W_2^2(\mu_t, \nu_t) \leq \exp(-2\alpha t) W_2^2(\mu_0, \nu_0)$$

using the following steps. First, compute the derivative of $t \mapsto W_2^2(\mu_t, \nu_t)$ using [Theorem 1.4.6](#). Next, apply the strong convexity inequality ([1.4.7](#)) to obtain two inequalities $\mathcal{F}(\mu_t) \geq \mathcal{F}(\nu_t) + \dots$ and $\mathcal{F}(\nu_t) \geq \mathcal{F}(\mu_t) + \dots$. Adding these two inequalities, deduce that $\partial_t W_2^2(\mu_t, \nu_t) \leq -2\alpha W_2^2(\mu_t, \nu_t)$.

▷ **Exercise 1.18** (Sharp KL convergence)

Let $(\mu_t)_{t \geq 0}$ denote the marginal law of the Langevin diffusion ([1.0.1](#)) with $\nabla^2 V \geq \alpha I_d \geq 0$.

- 1 Either prove via direct calculation that $\text{FI}(\mu_t \parallel \pi) \leq \exp(-2\alpha t) \text{FI}(\mu_0 \parallel \pi)$, or refer to the proof of [Theorem 1.2.30](#) in Chapter 2.
- 2 Consider the Lyapunov function $\mathcal{L}_t := A_t \text{FI}(\mu_t \parallel \pi) + 2B_t \text{KL}(\mu_t \parallel \pi) + W_2^2(\mu_t, \pi)$. Choose A_t, B_t carefully to ensure that $\dot{\mathcal{L}}_t \leq -\alpha \mathcal{L}_t$, and thereby deduce the bounds

$$\text{FI}(\mu_t \parallel \pi) \leq \frac{\alpha^2 W_2^2(\mu_0, \pi)}{\exp(2\alpha t) (1 - \exp(-\alpha t))^2}, \quad \text{KL}(\mu_t \parallel \pi) \leq \frac{\alpha W_2^2(\mu_0, \pi)}{2(\exp(2\alpha t) - 1)}.$$

Overview of the Convergence Results

▷ **Exercise 1.19** (Divergences at initialization)

Prove [Lemma 1.5.2](#). In fact, prove the following stronger fact: $\log \sup(\mu_0/\pi) \leq \frac{d}{2} \log \kappa$.

▷ **Exercise 1.20** (First moment of an exponential distribution)

Consider the probability measure $\pi \propto \exp(-\|\cdot\|)$ over $\mathbb{R}^d$. Show that the first moment behaves as $\int \|\cdot\| d\pi = d$. This yields a natural example of a log-concave distribution for which the Wasserstein distance at initialization is expected to scale as d , rather than $\sqrt{d}$.

▷ **Exercise 1.21** (Sharpness of the initialization bounds)

How sharp are [Lemma 1.5.2](#) and [Lemma 1.5.3](#)? Consider the case where π is a Gaussian.

In this chapter, we explore the connection between functional inequalities, such as the Poincaré and log-Sobolev inequalities, and the concentration of measure phenomenon.

2.1 Overview of the Inequalities

2.1.1 Relationships between the Inequalities

Let $C_c^\infty(\mathbb{R}^d)$ denote the set of compactly supported and smooth functions $\mathbb{R}^d \rightarrow \mathbb{R}$. The main inequalities that we study in this chapter are the following:

- the **Poincaré inequality (PI)**, as specialized to the Langevin diffusion (see [Example 1.2.23](#)):

$$\mathrm{var}_\pi(f) \leq C_{\mathrm{PI}} \mathbb{E}_\pi[\|\nabla f\|^2], \quad \text{for all } f \in C_c^\infty(\mathbb{R}^d).$$

- the **log-Sobolev inequality (LSI)**, as specialized to the Langevin diffusion (see [Example 1.2.27](#)):

$$\mathrm{ent}_\pi(f^2) \leq 2C_{\mathrm{LSI}} \mathbb{E}_\pi[\|\nabla f\|^2], \quad \text{for all } f \in C_c^\infty(\mathbb{R}^d).$$

- and **Talagrand's T_2 inequality**

$$\mathrm{KL}(\mu \parallel \pi) \geq \frac{1}{2C_{T_2}} W_2^2(\mu, \pi), \quad \text{for all } \mu \in \mathcal{P}_2(\mathbb{R}^d).$$

In addition, using the W_1 metric introduced in [Exercise 1.10](#), we consider

- **Talagrand's T_1 inequality**

$$\mathrm{KL}(\mu \parallel \pi) \geq \frac{1}{2C_{T_1}} W_1^2(\mu, \pi), \quad \text{for all } \mu \in \mathcal{P}_1(\mathbb{R}^d).$$

In many cases, arguments involving Poincaré and log-Sobolev inequalities hold more generally in the context of reversible Markov processes, and when this is the case we will try to use notation from Markov semigroup theory (e.g., writing $\mathbb{E}_\pi \Gamma(f)$ or $\mathcal{E}(f)$ instead of $\mathbb{E}_\pi[\|\nabla f\|^2]$) to indicate that this

is the case. However, for clarity of exposition, we do not dwell on this point, and we urge new readers to focus on the case in which the Markov process is the Langevin diffusion.

Although the Poincaré and log-Sobolev inequalities are stated above for compactly supported and smooth functions, once established they can be extended to a wider class of functions by density arguments; see [Bakry et al. \(2014, Chapter 3\)](#). Throughout the chapter, we omit mention of such approximation arguments.

Write $\text{PI}(C)$ to denote that the Poincaré inequality holds with constant C , and similarly for the other inequalities. We have the following relationships.

- The Bakry–Émery theorem ([Theorem 1.2.30](#)) shows that α -strong log-concavity of π implies that π satisfies $\text{LSI}(\alpha^{-1})$.
- The Otto–Villani theorem ([Exercise 1.16](#)) shows that $\text{LSI}(C)$ implies $T_2(C)$.
- Since $W_1 \leq W_2$, then $T_2(C)$ obviously implies $T_1(C)$. On the other hand, we will show below that $T_2(C)$ implies $\text{PI}(C)$ as well. Combined with the previous point, this shows that $\text{LSI}(C)$ implies $\text{PI}(C)$, which was shown directly in [Exercise 1.7](#).
- In general, PI and T_1 are incomparable.

2.1.2 Linearization of Transport Inequalities

To prove that $T_2(C)$ implies $\text{PI}(C)$, we linearize the transport cost. It will be convenient for future purposes to prove a more general version of the linearization principle.

Proposition 2.1.1 (Linearization of transport cost). *Let $c : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}_+$ be a lower semicontinuous cost function. Assume that $c(x, x) = 0$ for all $x \in \mathbb{R}^d$, that there exists $\delta > 0$ for which $c(x, y) \geq \delta \|x - y\|^2$ for all $x, y \in \mathbb{R}^d$, and that there is a measurable mapping $x \mapsto H_x > 0$ such that for each compact $K \subseteq \mathbb{R}^d$,*

$$\sup_{x \in K} \left| c(x + h, x) - \frac{1}{2} \langle h, H_x h \rangle \right| = o(\|h\|^2) \quad \text{as } h \rightarrow 0.$$

Then, for any $\mu \in \mathcal{P}(\mathbb{R}^d)$ and $f \in C^\infty(\mathbb{R}^d)$ which is bounded, whose gradient has compact support, and $\int f d\mu = 0$, it holds that

$$\liminf_{\varepsilon \searrow 0} \frac{1}{\varepsilon^2} \mathcal{T}_c(\mu, (1 + \varepsilon f) \mu) \geq \frac{(\int f^2 d\mu)^2}{2 \int \langle \nabla f(x), H_x^{-1} \nabla f(x) \rangle \mu(dx)}.$$

Proof sketch Fix $\lambda \in \mathbb{R}$. Using the dual formulation given in [Exercise 1.10](#),

$$\mathcal{T}_c(\mu, \nu) \geq \int \lambda \varepsilon f d\mu + \int (\lambda \varepsilon f)^c d\nu = \int (\lambda \varepsilon f)^c d\nu.$$

Here, $(\lambda \varepsilon f)^c(x) = \inf_{h \in \mathbb{R}^d} \{c(x + h, x) - \lambda \varepsilon f(x + h)\}$, and using the assumption on c and the compact support of f , one can justify that the infimum is attained at a point h with $\|h\| = O(\varepsilon)$. Then,

$$\begin{aligned} (\lambda \varepsilon f)^c(x) &= \inf_{h \in \mathbb{R}^d} \left\{ \frac{1}{2} \langle h, H_x h \rangle - \lambda \varepsilon f(x) - \lambda \varepsilon \langle \nabla f(x), h \rangle \right\} + o(\varepsilon^2) \\ &\geq -\lambda \varepsilon f(x) - \frac{\lambda^2 \varepsilon^2}{2} \langle \nabla f(x), H_x^{-1} \nabla f(x) \rangle + o(\varepsilon^2). \end{aligned}$$

Hence,

$$\mathcal{T}_c(\mu, (1 + \varepsilon f) \mu) \geq -\lambda \varepsilon^2 \int f^2 d\mu - \frac{\lambda^2 \varepsilon^2}{2} \int \langle \nabla f(x), H_x^{-1} \nabla f(x) \rangle \mu(dx)$$

and the result follows by optimizing over λ . $\square$

Corollary 2.1.2 (T_2 implies PI). *If π satisfies $T_2(C)$, then it satisfies PI(C).*

Proof Let $f \in C^\infty(\mathbb{R}^d)$ satisfy the assumption in the preceding proposition and apply the linearization to the quadratic cost $c(x, y) = \frac{1}{2} \|x - y\|^2$ with $H_x = I_d$ for all $x \in \mathbb{R}^d$. Then, $T_2(C)$ yields

$$2C \text{KL}((1 + \varepsilon f) \pi \parallel \pi) \geq W_2^2(\pi, (1 + \varepsilon f) \pi) \geq \frac{\varepsilon^2 (\int f^2 d\pi)^2}{\int \|\nabla f\|^2 d\pi} + o(\varepsilon^2).$$

On the other hand, the linearization (1.E.1) of the KL divergence in Exercise 1.7 yields

$$\text{KL}((1 + \varepsilon f) \pi \parallel \pi) = \frac{\varepsilon^2}{2} \int f^2 d\pi + o(\varepsilon^2).$$

Comparing terms proves the result. $\square$

In Exercise 2.1, we explore a perhaps more intuitive approach to linearizing the 2-Wasserstein distance via the Monge–Ampère equation.

2.2 Proofs via Markov Semigroup Theory

2.2.1 Commutation and Curvature

In Section 1.2.3, we introduced the iterated carré du champ operator Γ_2 , as well as the curvature-dimension condition $\Gamma_2(f) \geq \alpha \Gamma(f)$ (denoted $\text{CD}(\alpha, \infty)$). Here and throughout, we write $\Gamma(f) := \Gamma(f, f)$, $\Gamma_2(f) := \Gamma_2(f, f)$, and $\mathcal{E}(f) := \mathcal{E}(f, f)$ as a convenient shorthand. Since this condition plays a key role in the subsequent calculations, our goal is to demystify this idea.

By definition, the iterated carré du champ is

$$\Gamma_2(f) = \frac{1}{2} \{ \mathcal{L}\Gamma(f) - 2\Gamma(f, \mathcal{L}f) \}.$$

For the case of the Langevin diffusion with carré du champ $\Gamma(f) = \|\nabla f\|^2$,

$$\Gamma_2(f) = \frac{1}{2} \{ \mathcal{L}(\|\nabla f\|^2) - 2\langle \nabla f, \nabla \mathcal{L}f \rangle \}. \quad (2.2.1)$$

Recall that $\mathcal{L}f = \Delta f - \langle \nabla V, \nabla f \rangle$, where V is the potential. Let us begin with the simple case in which $V = 0$, so $\mathcal{L}$ is the Laplacian Δ (the generator of $\sqrt{2} B$, where B is standard Brownian motion). In this case, the iterated carré du champ turns out to simply be the operator $\Gamma_2(f) = \|\nabla^2 f\|_{\text{HS}}^2$, which is known as the **Bochner identity**:

$$\frac{1}{2} \Delta(\|\nabla f\|^2) = \langle \nabla \Delta f, \nabla f \rangle + \|\nabla^2 f\|_{\text{HS}}^2. \quad (2.2.2)$$

Consequently, Δ satisfies $\text{CD}(0, \infty)$.

It may seem strange at first sight to give such a fancy name to the seemingly innocuous identity (2.2.2),

which is a simple exercise in calculus. However, the importance of the Bochner identity begins to reveal itself through the following fact: the identity continues to make sense on a Riemannian manifold, except that there is an extra term involving the *Ricci curvature* of the manifold.

$$\frac{1}{2} \Delta(\|\nabla f\|^2) = \langle \nabla \Delta f, \nabla f \rangle + \|\nabla^2 f\|_{\text{HS}}^2 + \text{Ric}(\nabla f, \nabla f) .$$

We will defer a fuller discussion of Riemannian geometry for later, but for now we can get a hint at the role of the curvature by observing that the Bochner identity (2.2.2) on $\mathbb{R}^d$ follows from the equation¹

$$\nabla \Delta f - \Delta \nabla f = 0 \quad (2.2.3)$$

by taking the inner product with ∇f and applying the identity

$$\frac{1}{2} \Delta(\|\nabla f\|^2) = \text{div}(\nabla^2 f \nabla f) = \langle \nabla \Delta f, \nabla f \rangle + \|\nabla^2 f\|_{\text{HS}}^2 .$$

In turn, the equation (2.2.3) shows that the Laplacian commutes with the gradient operator, which is true because partial derivatives commute on $\mathbb{R}^d$; this is a manifestation of the fact that $\mathbb{R}^d$ is *flat*. In contrast, the very definition of curvature on a Riemannian manifold is usually based upon the *lack* of commutativity of differential operators.²

Turning now to the Langevin generator $\mathcal{L}$, the identity (2.2.3) is replaced by

$$\nabla \mathcal{L} f - \mathcal{L} \nabla f = -\nabla^2 V \nabla f . \quad (2.2.4)$$

Hence, the commutator of ∇ and $\mathcal{L}$ brings out the curvature of the measure $\pi \propto \exp(-V)$, and the plan is to exploit this in order to prove functional inequalities. The identity (2.2.4) then yields the following formula for the iterated carré du champ:

$$\Gamma_2(f) = \|\nabla^2 f\|_{\text{HS}}^2 + \langle \nabla f, \nabla^2 V \nabla f \rangle . \quad (2.2.5)$$

In particular, if $\nabla^2 V \geq \alpha I_d > 0$, then the curvature-dimension condition $\text{CD}(\alpha, \infty)$ holds, which was asserted as [Theorem 1.2.31](#).

As a first illustration of the use of curvature, let us generalize the calculation in [Exercise 1.8](#) from the Ornstein–Uhlenbeck diffusion to the Langevin diffusion. From the computation $\partial_t \text{var}_\pi P_t f = -2\mathcal{E}(P_t f) = -2 \int \Gamma(P_t f) d\pi$ and $\text{var}_\pi P_t f \rightarrow 0$ as $t \rightarrow \infty$ by ergodicity, we obtain the identity $\text{var}_\pi f = 2 \int_0^\infty \mathcal{E}(P_t f) dt$. For the Ornstein–Uhlenbeck semigroup, we have the explicit identity $\nabla P_t f = \exp(-t) P_t \nabla f$, which leads to $\mathcal{E}(P_t f) \leq \exp(-2t) \mathcal{E}(f)$ and hence $\text{var}_\pi f \leq \mathcal{E}(f)$, which is a Poincaré inequality for the standard Gaussian measure.

To extend this idea to more general diffusions, we show that the curvature-dimension condition is equivalent to a gradient bound.

Theorem 2.2.6 (Gradient bound). *The following are equivalent.*

- 1 *The curvature-dimension condition $\text{CD}(\alpha, \infty)$ holds.*
- 2 *For all f and $t \geq 0$,*

$$\Gamma(P_t f) \leq \exp(-2\alpha t) P_t \Gamma(f) .$$

¹ Here, Δ acts on ∇u component by component.

² Loosely speaking, the idea of curvature is that travelling in direction u and then direction v is not exactly the same as travelling in direction v and direction u . Algebraically, this is captured by studying the difference between differentiating along vector field X and then vector field Y , or vice versa.

Proof (1) $\implies$ (2): We differentiate the following quantity:

$$\partial_s [P_s \Gamma(P_{t-s} f)] = P_s (\mathcal{L} \Gamma(P_{t-s} f) - 2 \Gamma(P_{t-s} f, \mathcal{L} P_{t-s} f)).$$

Applying the definition of the iterated carré du champ and $\text{CD}(\alpha, \infty)$ yields

$$\partial_s [P_s \Gamma(P_{t-s} f)] = 2 P_s \Gamma_2(P_{t-s} f) \geq 2 \alpha P_s \Gamma(P_{t-s} f).$$

Integrating this from $s = 0$ to $s = t$ yields $P_t \Gamma(f) \geq \exp(2\alpha t) \Gamma(P_t f)$.

(2) $\implies$ (1): A Taylor expansion of the inequality around $t = 0$ yields

$$\Gamma(f) + 2t \Gamma(f, \mathcal{L} f) + o(t) \leq \Gamma(f) + t \{-2\alpha \Gamma(f) + \mathcal{L} \Gamma(f)\} + o(t),$$

which clearly recovers the $\text{CD}(\alpha, \infty)$ condition $\Gamma_2 \geq \alpha \Gamma$. $\square$

Substituting this into $\text{var}_\pi f = 2 \int_0^\infty \mathcal{E}(P_t f) dt$ now shows that as soon as the Markov semigroup satisfies $\text{CD}(\alpha, \infty)$, the Poincaré inequality $\text{PI}(\alpha^{-1})$ holds.

2.2.2 The Brascamp–Lieb Inequality

In this section, we prove a far-reaching refinement of the preceding argument and thereby establish the Brascamp–Lieb inequality, a strong form of the Poincaré inequality. This inequality will also gain a natural interpretation via a diffusion process in Section 10.3.

The proof method in this section is known as **Hörmander’s L^2 method**. The starting point is to write down a dual form of the Poincaré inequality.³

Lemma 2.2.7 (Dualizing the Poincaré inequality). *Let $\pi \propto \exp(-V)$ be a probability measure on $\mathbb{R}^d$, where V is continuously differentiable; let $\mathcal{L}$ be the corresponding Langevin generator. Suppose that $A : \mathbb{R}^d \rightarrow \text{PD}(d)$ is a matrix-valued function mapping into the space of symmetric positive definite matrices such that for all smooth $u : \mathbb{R}^d \rightarrow \mathbb{R}$,*

$$\mathbb{E}_\pi[(\mathcal{L} u)^2] \geq \mathbb{E}_\pi\langle \nabla u, A \nabla u \rangle. \quad (2.2.8)$$

Then, for all smooth $f : \mathbb{R}^d \rightarrow \mathbb{R}$ it holds that

$$\text{var}_\pi f \leq \mathbb{E}_\pi\langle \nabla f, A^{-1} \nabla f \rangle.$$

Proof Assume that $\mathbb{E}_\pi f = 0$. Recall that $\mathbb{E}_\pi \mathcal{L} u = 0$ for any u , so that $\mathbb{E}_\pi f = 0$ is a *necessary* condition for the solvability of the Poisson equation $-\mathcal{L} u = f$. For simplicity, we will assume that this condition is also sufficient.

If we express $\mathbb{E}_\pi[f^2]$ in terms of u and apply integration by parts (Theorem 1.2.14) and Cauchy–Schwarz, we obtain

$$\mathbb{E}_\pi[f^2] = \mathbb{E}_\pi[f(-\mathcal{L} u)] = \mathbb{E}_\pi\langle \nabla f, \nabla u \rangle \leq \sqrt{\mathbb{E}_\pi\langle \nabla f, A^{-1} \nabla f \rangle \mathbb{E}_\pi\langle \nabla u, A \nabla u \rangle}.$$

By assumption, $\mathbb{E}_\pi\langle \nabla u, A \nabla u \rangle \leq \mathbb{E}_\pi[(\mathcal{L} u)^2] = \mathbb{E}_\pi[f^2]$, which establishes the result. $\square$

³The idea of dualizing the Poincaré inequality also appears in Exercise 2.1, in which the Poincaré inequality is deduced from an inequality on $(-\mathcal{L})^{-1}$.

The point is that the condition (2.2.8) can now be checked with the help of curvature. Suppose that $\pi \propto \exp(-V)$ where V is twice continuously differentiable and strictly convex. Then, using integration by parts (Theorem 1.2.14),

$$\begin{aligned} \mathbb{E}_\pi[(\mathcal{L}u)^2] &= -\mathbb{E}_\pi\langle \nabla u, \nabla \mathcal{L}u \rangle = \mathbb{E}_\pi \left[\Gamma_2(u) - \underbrace{\frac{1}{2} \mathcal{L}(\|\nabla u\|^2)}_{\text{by (2.2.1)}} \right] = \underbrace{\mathbb{E}_\pi \Gamma_2(u)}_{\text{because } \mathbb{E}_\pi \mathcal{L}=0} \\ &= \mathbb{E}_\pi \left[\underbrace{\|\nabla^2 u\|_{\text{HS}}^2 + \langle \nabla u, \nabla^2 V \nabla u \rangle}_{\text{by (2.2.5)}} \right]. \end{aligned}$$

Applying the lemma, we obtain the following result.

Theorem 2.2.9 (Bascamp–Lieb inequality). *Let $\pi \propto \exp(-V)$, where V is strictly convex on $\mathbb{R}^d$ and twice continuously differentiable. Then, for all $f : \mathbb{R}^d \rightarrow \mathbb{R}$,*

$$\text{var}_\pi f \leq \mathbb{E}_\pi \langle \nabla f, (\nabla^2 V)^{-1} \nabla f \rangle.$$

When $\nabla^2 V \geq \alpha I_d > 0$, then this implies that a Poincaré inequality holds for π with constant $C_{\text{PI}} \leq 1/\alpha$. However, the Bascamp–Lieb inequality is much stronger, as it allows us to take advantage of non-uniform convexity.

Remark 2.2.10. Lemma 2.2.7 shows that $\text{PI}(\alpha^{-1})$ is equivalent to the *integrated* version of the $\text{CD}(\alpha, \infty)$ condition: $\int \Gamma_2(f) d\pi \geq \alpha \int \Gamma(f) d\pi$ for all f .

In Exercise 2.3, we give another proof of Theorem 2.2.9 by linearizing a transport inequality. First, we introduce the transport cost.

Definition 2.2.11. The **Bregman transport cost** for the potential V , denoted $\mathcal{D}_V$, is the transport cost associated with the **Bregman divergence**

$$D_V(x, y) := V(x) - V(y) - \langle \nabla V(y), x - y \rangle,$$

i.e., we set

$$\mathcal{D}_V(\mu, \nu) := \inf_{\gamma \in \mathcal{C}(\mu, \nu)} \int D_V(x, y) \gamma(dx, dy).$$

The Bregman transport cost will also play a key role in Section 10.3, in which we will prove the following transport inequality (see also Exercise 2.3).

Theorem 2.2.12 (Bregman transport inequality). *Suppose that $V : \mathbb{R}^d \rightarrow \mathbb{R}$ is continuously differentiable. Then, for $\pi \propto \exp(-V)$ and all $\mu \in \mathcal{P}(\mathbb{R}^d)$,*

$$\mathcal{D}_V(\mu, \pi) \leq \text{KL}(\mu \parallel \pi).$$

Actually, convexity of V is not necessary for the theorem to hold, although the Bregman transport cost $\mathcal{D}_V$ is only guaranteed to be non-negative when V is convex. When V is strongly convex, $\nabla^2 V \geq \alpha I_d > 0$, then $D_V(x, y) \geq \frac{\alpha}{2} \|x - y\|^2$, so the Bregman transport inequality implies $\text{T}_2(\alpha^{-1})$.

2.2.3 Proof of the Bakry–Émery Theorem

In this section, we prove the Bakry–Émery theorem ([Theorem 1.2.30](#)), i.e., that $\text{CD}(\alpha, \infty)$ implies $\text{LSI}(\alpha^{-1})$. We introduce the notion of a diffusion semigroup.

Definition 2.2.13. The Markov semigroup $(P_t)_{t \geq 0}$ is a **diffusion semigroup** if for all functions $f, g \in L^2(\pi)$ in the domain of the carré du champ Γ and all $\phi : \mathbb{R} \rightarrow \mathbb{R}$, the chain rule holds:

$$\Gamma(\phi \circ f, g) = \phi'(f) \Gamma(f, g). \quad (2.2.14)$$

More generally, for functions $f_1, \dots, f_k$ and $\Psi : \mathbb{R}^k \rightarrow \mathbb{R}$,

$$\Gamma(\Psi(f_1, \dots, f_k), g) = \sum_{i=1}^k (\partial_i \Psi)(f_1, \dots, f_k) \Gamma(f_i, g).$$

Equivalently, the chain rule can be stated for the generator ([Exercise 2.8](#)):

$$\mathcal{L}\phi(f) = \phi'(f) \mathcal{L}f + \phi''(f) \Gamma(f). \quad (2.2.15)$$

The chain rule is satisfied for the Langevin diffusion whose carré du champ is given by $\Gamma(f, g) = \langle \nabla f, \nabla g \rangle$, and more generally this assumption encodes the fact that the Markov process is a diffusion. Since we are mainly interested in diffusion processes, this is not a restrictive assumption, but it indicates that the following proof will fail for Markov processes on discrete state spaces.

Proof of the Bakry–Émery theorem ([Theorem 1.2.30](#)) Given a smooth positive function f and a function $\phi : \mathbb{R}_+ \rightarrow \mathbb{R}$, we differentiate $t \mapsto \int \phi(P_t f) d\pi$. We are primarily interested in the case $\phi(x) := x \log x$, but carrying out the calculation for a general ϕ clarifies the structure of the argument. Using the Markov semigroup calculus,

$$\partial_t \int \phi(P_t f) d\pi = \int \phi'(P_t f) \mathcal{L}P_t f d\pi = -\mathcal{E}(\phi' \circ P_t f, P_t f).$$

This yields the representation

$$\begin{aligned} \text{ent}_\pi^\phi f &:= \int \phi(f) d\pi - \phi\left(\int f d\pi\right) = -\int_0^\infty \left(\partial_t \int \phi(P_t f) d\pi\right) dt \\ &= \int_0^\infty \mathcal{E}(\phi' \circ P_t f, P_t f) dt. \end{aligned}$$

We now specialize our calculations to the entropy function $\phi(x) = x \log x$ and use reversibility of the semigroup and the chain rule.

$$\mathcal{E}(\phi' \circ P_t f, P_t f) = \mathcal{E}(\log P_t f, P_t f) = \int P_t f \Gamma(\log P_t f) d\pi.$$

Since $\partial_t \log P_t f = (\mathcal{L}P_t f)/P_t f = \mathcal{L} \log P_t f + \Gamma(\log P_t f)$ by (2.2.15),

$$\partial_t \mathcal{E}(\log P_t f, P_t f) = \int \{\mathcal{L}P_t f \Gamma(\log P_t f) + 2P_t f \Gamma(\log P_t f, \mathcal{L} \log P_t f + \Gamma(\log P_t f))\} d\pi.$$

Furthermore, $\int \mathcal{L} P_t f \Gamma(\log P_t f) d\pi = - \int \Gamma(P_t f, \Gamma(\log P_t f)) d\pi$, so

$$\begin{aligned} \partial_t \mathcal{E}(\log P_t f, P_t f) &= \int \{-\mathcal{L} P_t f \Gamma(\log P_t f) + 2P_t f \Gamma(\log P_t f, \mathcal{L} \log P_t f)\} d\pi \\ &= - \int P_t f \Gamma_2(\log P_t f) d\pi. \end{aligned}$$

Under $\text{CD}(\alpha, \infty)$, this readily yields

$$\mathcal{E}(\log P_t f, P_t f) \leq \exp(-2\alpha t) \mathcal{E}(\log f, f). \quad (2.2.16)$$

In words: under $\text{CD}(\alpha, \infty)$, the Fisher information (introduced in [Example 1.2.27](#)) decays exponentially fast. Another proof of this fact is given as [Exercise 3.5](#).

Substituting this into the representation above,

$$\text{ent}_\pi f = \int_0^\infty \mathcal{E}(\log P_t f, P_t f) dt \leq \mathcal{E}(\log f, f) \int_0^\infty \exp(-2\alpha t) dt \leq \frac{1}{2\alpha} \mathcal{E}(\log f, f),$$

which is the log-Sobolev inequality. $\square$

2.2.4 Local Inequalities

The curvature-dimension condition is a *pointwise* inequality that does not explicitly involve the stationary distribution π . Interestingly, this means that the Poincaré and log-Sobolev inequalities can be strengthened to local inequalities, which turn out to be equivalent to the curvature-dimension condition. To illustrate this, we state local versions of the Poincaré inequality, along with an interesting reverse inequality.

Theorem 2.2.17 (Local Poincaré inequality). *Let $(P_t)_{t \geq 0}$ be a diffusion semigroup. The following are equivalent.*

- 1 The curvature-dimension condition $\text{CD}(\alpha, \infty)$ holds.
- 2 For all f and $t \geq 0$,

$$P_t(f^2) - (P_t f)^2 \leq \frac{1 - \exp(-2\alpha t)}{\alpha} P_t \Gamma(f). \quad (2.2.18)$$

- 3 For all f and $t \geq 0$,

$$P_t(f^2) - (P_t f)^2 \geq \frac{\exp(2\alpha t) - 1}{\alpha} \Gamma(P_t f). \quad (2.2.19)$$

Consider the local Poincaré inequality (2.2.18). When evaluated at a point x , and keeping in mind that $P_t f(x) = \mathbb{E}_{\delta_x P_t} f$ where $\delta_x P_t$ is the law of the Markov process started at δ_x , it can be interpreted as follows: *for all x and all $t \geq 0$, the measure $\delta_x P_t$ satisfies a Poincaré inequality with constant $(1 - \exp(-2\alpha t))/\alpha$. If $\alpha > 0$, then as $t \rightarrow \infty$ it recovers the Poincaré inequality for the stationary measure with constant $1/\alpha$, but the local inequality is considerably stronger: for example, it is meaningful even for $\alpha \leq 0$. When $\alpha = 0$, the constant should be interpreted in terms of its limiting value, i.e., $\lim_{\alpha \rightarrow 0} (1 - \exp(-2\alpha t))/\alpha = 2t$. Similarly, the inequality (2.2.19) can be understood as a local reverse Poincaré inequality.*

To give the flavor for the arguments, let ϕ be a convex function. The proofs begin with the following differentiation:

$$\partial_s P_s \phi(P_{t-s} f) = P_s (\mathcal{L} \phi(P_{t-s} f) - \phi'(P_{t-s} f) \mathcal{L} P_{t-s} f) = P_s (\phi''(P_{t-s} f) \Gamma(P_{t-s} f)),$$

where we applied the diffusion chain rule (2.2.15). It yields

$$P_t \phi(f) - \phi(P_t f) = \int_0^t P_s (\phi''(P_{t-s} f) \Gamma(P_{t-s} f)) ds.$$

By Jensen's inequality, the left-hand side is always non-negative, but the right-hand side quantifies this effect. Indeed, Markov semigroup calculations are designed to extract useful information from the lack of commutativity of the semigroup P_t and the application of the non-linear function ϕ . If we apply this to $\phi : x \mapsto x^2$, we obtain

$$P_t(f^2) - (P_t f)^2 = 2 \int_0^t P_s \Gamma(P_{t-s} f) ds. \quad (2.2.20)$$

Under $\text{CD}(\alpha, \infty)$, we have $P_s \Gamma(P_{t-s} f) \leq \exp(-2\alpha(t-s)) P_t \Gamma(f)$ from Theorem 2.2.6, and substituting this into (2.2.20) yields the local Poincaré inequality (2.2.18). The other implications in Theorem 2.2.17 are left as Exercise 2.7.

A salient feature of these arguments is that by their pointwise nature—i.e., they are not integrated w.r.t. π —we never invoke integration by parts. They are also perfectly valid even for an infinite invariant reference measure, e.g., Lebesgue measure for standard Brownian motion.

Finally, we also record the local log-Sobolev inequality (a proof of the implication $1 \Rightarrow 3$ for the Langevin diffusion is given in Exercise 4.4).

Theorem 2.2.21 (Local log-Sobolev inequality). *Let $(P_t)_{t \geq 0}$ be a diffusion semigroup. The following are equivalent.*

- 1 The curvature-dimension condition $\text{CD}(\alpha, \infty)$ holds.
- 2 For all f and $t \geq 0$, $\sqrt{\Gamma(P_t f)} \leq \exp(-\alpha t) P_t \sqrt{\Gamma(f)}$.
- 3 For all $f \geq 0$ and $t \geq 0$,

$$P_t(f \log f) - P_t f \log P_t f \leq \frac{1 - \exp(-2\alpha t)}{2\alpha} P_t \frac{\Gamma(f)}{f}.$$

- 4 For all $f \geq 0$ and $t \geq 0$,

$$P_t(f \log f) - P_t f \log P_t f \geq \frac{\exp(2\alpha t) - 1}{2\alpha} \frac{\Gamma(P_t f)}{P_t f}. \quad (2.2.22)$$

2.2.5 Convergence in Rényi Divergence

One curiosity is that the log-Sobolev inequality directly implies a Poincaré inequality (Exercise 1.7), and yet the convergence guarantees implied by these inequalities for the Langevin diffusion are incomparable, because they apply to different metrics (χ^2 vs. KL). It turns out that these convergence guarantees can be placed in the same framework by introducing the family of *Rényi divergences*. Rényi divergences have also gained importance in recent research due to applications to differential privacy (Mironov, 2017).

Definition 2.2.23. For $q > 1$, the **Rényi divergence** of order q between μ and π is defined by

$$\mathcal{R}_q(\mu \parallel \pi) := \frac{1}{q-1} \log \int \left(\frac{d\mu}{d\pi} \right)^q d\pi \quad (2.2.24)$$

if $\mu \ll \pi$, and $\mathcal{R}_q(\mu \parallel \pi) := +\infty$ otherwise.

Rényi divergences are monotonic in the order: if $1 < q \leq q' < \infty$, then $\mathcal{R}_q \leq \mathcal{R}_{q'}$ (this follows from Jensen's inequality, see [Exercise 2.9](#)). Rényi divergences can also be defined for $q < 1$, but these cases are less useful for the purpose of this book. Some notable special cases include:

- As $q \searrow 1$, $\mathcal{R}_q \searrow \text{KL}$; hence, we set $\mathcal{R}_1 = \text{KL}$.
- For $q = 2$, we have the identity $\mathcal{R}_2 = \log(1 + \chi^2)$.
- As $q \nearrow \infty$, $\mathcal{R}_q \nearrow \mathcal{R}_\infty$, where $\mathcal{R}_\infty(\mu \parallel \pi) := \log \left\| \frac{d\mu}{d\pi} \right\|_{L^\infty(\pi)}$.

Remarkably, [Vempala and Wibisono \(2019\)](#) show that the Poincaré and log-Sobolev inequalities imply convergence of the Langevin diffusion in every Rényi divergence.

Theorem 2.2.25 (Rényi convergence, [Vempala and Wibisono \(2019\)](#)). *Let $(P_t)_{t \geq 0}$ be a reversible diffusion Markov semigroup, and let $(\pi_t)_{t \geq 0}$ denote the law of the Markov process associated with the semigroup.*

1 *Suppose that a log-Sobolev inequality holds with constant C_{LSI} . Then, for all $q \geq 1$,*

$$\mathcal{R}_q(\pi_t \parallel \pi) \leq \exp\left(-\frac{2t}{qC_{\text{LSI}}}\right) \mathcal{R}_q(\pi_0 \parallel \pi).$$

2 *Suppose that a Poincaré inequality holds with constant C_{PI} . Then, for all $q \geq 2$,*

$$\exp \mathcal{R}_q(\pi_t \parallel \pi) - 1 \leq \exp\left(-\frac{4t}{qC_{\text{PI}}}\right) (\exp \mathcal{R}_q(\pi_0 \parallel \pi) - 1).$$

Proof We begin by differentiating the Rényi divergence in time. Let $\rho_t := \frac{d\pi_t}{d\pi} = P_t \rho_0$. Applying the chain rule for the carré du champ,

$$\begin{aligned} \partial_t \mathcal{R}_q(\pi_t \parallel \pi) &= \frac{1}{q-1} \frac{\partial_t \int \rho_t^q d\pi}{\int \rho_t^q d\pi} = \frac{q}{q-1} \frac{\int \rho_t^{q-1} \mathcal{L} \rho_t d\pi}{\int \rho_t^q d\pi} = -\frac{q}{q-1} \frac{\int \Gamma(\rho_t^{q-1}, \rho_t) d\pi}{\int \rho_t^q d\pi} \\ &= -\frac{4}{q} \frac{\mathcal{E}(\rho_t^{q/2})}{\int \rho_t^q d\pi}. \end{aligned}$$

Log-Sobolev case. The log-Sobolev inequality reads (due to the chain rule) as

$$2C_{\text{LSI}} \mathcal{E}(f) \geq \text{ent}_\pi(f^2).$$

Applying this to $f = \rho^{q/2}$, we obtain

$$\begin{aligned} 2C_{\text{LSI}} \mathcal{E}(\rho^{q/2}) &\geq q \int \rho^q \log \rho d\pi - \left(\int \rho^q d\pi \right) \log \left(\int \rho^q d\pi \right) \\ &= q \partial_q \int \rho^q d\pi - \left(\int \rho^q d\pi \right) \log \left(\int \rho^q d\pi \right) \end{aligned}$$

and hence

$$\begin{aligned}
\frac{4}{q} \frac{\mathcal{E}(\rho^{q/2})}{\int \rho^q d\pi} &\geq \frac{2}{C_{\text{LSI}}} \partial_q \log \int \rho^q d\pi - \frac{2}{q C_{\text{LSI}}} \log \int \rho^q d\pi \\
&= \frac{2}{C_{\text{LSI}}} \partial_q [(q-1) \mathcal{R}_q(\rho\pi \parallel \pi)] - \frac{2(q-1)}{q C_{\text{LSI}}} \mathcal{R}_q(\rho\pi \parallel \pi) \\
&= \frac{2}{C_{\text{LSI}}} \mathcal{R}_q(\rho\pi \parallel \pi) + \frac{2(q-1)}{C_{\text{LSI}}} \underbrace{\partial_q \mathcal{R}_q(\rho\pi \parallel \pi)}_{\geq 0} - \frac{2(q-1)}{q C_{\text{LSI}}} \mathcal{R}_q(\rho\pi \parallel \pi) \\
&\geq \frac{2}{q C_{\text{LSI}}} \mathcal{R}_q(\rho\pi \parallel \pi)
\end{aligned}$$

where we used the fact that the Rényi divergence is monotonic in the order.

Poincaré case. Next, applying a Poincaré inequality to $f = \rho^{q/2}$,

$$\begin{aligned}
C_{\text{PI}} \mathcal{E}(\rho^{q/2}) &\geq \text{var}_\pi(\rho^{q/2}) = \int \rho^q d\pi - \left(\int \rho^{q/2} d\pi \right)^2 \\
&= \left(\int \rho^q d\pi \right) \left[1 - \frac{\exp((q-2) \mathcal{R}_{q/2}(\rho\pi \parallel \pi))}{\exp((q-1) \mathcal{R}_q(\rho\pi \parallel \pi))} \right] \\
&\geq \left(\int \rho^q d\pi \right) \{1 - \exp(-\mathcal{R}_q(\rho\pi \parallel \pi))\}
\end{aligned}$$

where we used the monotonicity of the Rényi divergence in the order. Hence,

$$\frac{4}{q} \frac{\mathcal{E}(\rho^{q/2})}{\int \rho^q d\pi} \geq \frac{4}{q C_{\text{PI}}} \{1 - \exp(-\mathcal{R}_q(\rho\pi \parallel \pi))\}.$$

This is equivalent to $\partial_t(\exp \mathcal{R}_q(\pi_t \parallel \pi) - 1) \leq -(4/q C_{\text{PI}})(\exp \mathcal{R}_q(\pi_t \parallel \pi) - 1)$. $\square$

To interpret this theorem, the Poincaré result states that after an initial waiting period of time $O(q C_{\text{PI}} \mathcal{R}_q(\pi_0 \parallel \pi))$, the Rényi divergence starts decaying exponentially fast. On the other hand, the log-Sobolev inequality implies exponentially fast convergence from the outset. In particular, for $q = 2$, we see that whereas a Poincaré inequality implies exponential decay of χ^2 , a log-Sobolev inequality implies exponential decay of $\log(1 + \chi^2)$, which is substantially stronger.

Another approach to Rényi convergence under $\text{CD}(\alpha, \infty)$ is based on Harnack inequalities, see [Exercise 2.17](#).

2.2.6 Beyond Strong Log-Concavity

Most of the results thus far have assumed $\text{CD}(\alpha, \infty)$ for $\alpha > 0$, which in the case of the Langevin diffusion is equivalent to strong log-concavity of π . What can we assert about the validity of functional inequalities more generally?

First, consider the case $\alpha = 0$: π is log-concave but not uniformly so. For example, π could be the Laplace distribution with density $\pi \propto \exp(-\frac{1}{2}|\cdot|)$ on $\mathbb{R}$. In this case, since π only has subexponential tails, we cannot hope for a log-Sobolev inequality due to [Theorem 2.4.3](#) below.⁴ The Brascamp–Lieb inequality ([Theorem 2.2.9](#)) holds as a useful substitute for the Poincaré inequality, but a deep conjecture

⁴However, if we further assume that π has support with diameter R , then a log-Sobolev inequality holds for π with constant $O(R^2)$; see [Exercise 3.9](#).

of Kannan, Lovász, and Simonovits (Kannan et al., 1995) goes much further: they conjectured that any log-concave measure π on $\mathbb{R}^d$ which is isotropic (i.e., if $X \sim \pi$ then $\text{cov } X = I_d$) satisfies a Poincaré inequality with a dimension-free constant $C_{\text{PI}} \lesssim 1$. This is known as the **Kannan–Lovász–Simonovits (KLS) conjecture**. By considering linear test functions of the form $x \mapsto \langle a, x \rangle$, one has $C_{\text{PI}} \geq 1$, so the conjecture asserts that linear functions nearly saturate the spectral gap inequality for log-concave measures. The KLS conjecture has inspired a considerable amount of research (including Theorem 2.4.22), see Guédon and Milman (2011); Eldan (2013); Lee and Vempala (2017); Chen (2021); Klartag and Lehec (2022), culminating in the current state-of-the-art result:

Theorem 2.2.26 (Klartag (2023)). *If π is an isotropic log-concave measure over $\mathbb{R}^d$, then it satisfies a Poincaré inequality with $C_{\text{PI}} \lesssim \log d$.*

In the non-convex case, the Poincaré and log-Sobolev inequalities can still hold, but with constants that scale exponentially with the amount of non-convexity. The most favorable case is when the potential V has a benign landscape. In this setting, there is an explicit expression for the low-temperature limit of the Poincaré and log-Sobolev constants, which shows that they scale as β^{-1} (where β is the inverse temperature).

Theorem 2.2.27 (Chewi and Stromme (2024)). *Assume that V is C^2 , admits a unique minimizer $x_\star$, and that there exists $C > 0$ such that $\Delta V \leq C(1 + \|\nabla V\|^2)$. Furthermore, assume that V satisfies a PL inequality with constant $C_{\text{PL}} < \infty$, that is,*

$$V - \inf V \leq \frac{C_{\text{PL}}}{2} \|\nabla V\|^2.$$

For $\beta > 0$, let π_β be the probability measure with density $\pi_\beta \propto \exp(-\beta V)$. Then,

$$\lim_{\beta \rightarrow \infty} \beta C_{\text{PI}}(\pi_\beta) = \frac{1}{\lambda_{\min}(\nabla^2 V(x_\star))}, \quad \lim_{\beta \rightarrow \infty} \beta C_{\text{LSI}}(\pi_\beta) = C_{\text{PL}}.$$

Note that the second formula connects the PL constant of V to the asymptotic limit of the LSI constant of π_β , and that the LSI is the analogue of the PL condition for the KL divergence over the Wasserstein space (see Section 1.4.2). This is somewhat analogous to Theorem 1.4.5, which connects the strong convexity of V to strong convexity of the KL divergence over the Wasserstein space, but the connection between the PL conditions is only true in the low-temperature limit. Interestingly, the asymptotic limit of the Poincaré constant only depends on the local behavior of V via the curvature at the minimizer, although Theorem 2.2.27 does impose global assumptions (that V admits a unique minimizer and satisfies the PL condition).

When V violates the unique minimizer assumption, the functional inequality constants generally scale exponentially. A special case of the **Eyring–Kramers formula** (Eyring, 1935; Kramers, 1940) asserts that if V has two separated minima x_0, x_1 , then the Poincaré constant for π scales exponentially in the size of the “energy barrier”. See Menz and Schlichting (2014) for a precise statement.

The next section is devoted to techniques for establishing functional inequalities in various settings.

2.3 Operations Preserving Functional Inequalities

To further expand the class of distributions known to satisfy functional inequalities, we will show in this section that functional inequalities are stable under various common operations on probability measures. We let $C_{\text{PI}}(\pi)$ denote the Poincaré constant of a probability measure π , and similarly write $C_{\text{LSI}}(\pi)$, $C_{\text{T}_2}(\pi)$, etc.

2.3.1 Bounded Perturbation

Suppose that π satisfies a functional inequality, and that μ is another probability measure satisfying $0 < c \leq \frac{d\mu}{d\pi} \leq C < \infty$. Then, it often follows that μ also satisfies the same functional inequality, with a worse constant. This furnishes a large class of examples of non-log-concave measures satisfying functional inequalities.

Proposition 2.3.1 (Holley–Stroock perturbation, [Holley and Stroock \(1987\)](#)). *Suppose that π satisfies either a Poincaré or log-Sobolev inequality. Then, if μ is such that $\frac{d\mu}{d\pi} \propto \exp B$, where $\text{osc } B := \sup B - \inf B \leq b$, then μ also satisfies the corresponding functional inequality with*

$$C_{\text{PI}}(\mu) \leq C_{\text{PI}}(\pi) \exp b \quad \text{or} \quad C_{\text{LSI}}(\mu) \leq C_{\text{LSI}}(\pi) \exp b$$

respectively.

Proof The key is to find a variational principle for the variance or for the entropy. For the variance, for any $v \in \mathcal{P}(\mathbb{R}^d)$ and $f : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\text{var}_v f = \inf_{m \in \mathbb{R}} \mathbb{E}_v[|f - m|^2].$$

From this, if Z denotes the normalizing constant,

$$\begin{aligned} \text{var}_\mu f &= \inf_{m \in \mathbb{R}} \int |f - m|^2 \frac{\exp B}{Z} d\pi \leq \frac{\exp \sup B}{Z} \inf_{m \in \mathbb{R}} \int |f - m|^2 d\pi = \frac{\exp \sup B}{Z} \text{var}_\pi f \\ &\leq \frac{C_{\text{PI}}(\pi) \exp \sup B}{Z} \int \Gamma(f) d\pi \leq C_{\text{PI}}(\pi) \exp(\sup B - \inf B) \int \Gamma(f) \frac{\exp B}{Z} d\pi \\ &= C_{\text{PI}}(\pi) \exp b \int \Gamma(f) d\mu. \end{aligned}$$

The proof for the log-Sobolev inequality is similar once we have the variational principle

$$\text{ent}_v f = \inf_{t > 0} \mathbb{E}_v \left[\underbrace{f \log \frac{f}{t} - f + t}_{\geq 0} \right]$$

for any $f : \mathbb{R}^d \rightarrow \mathbb{R}_+$, which we leave as [Exercise 2.12](#). □

One can also state a perturbation principle for the T_2 inequality ([Gozlan et al., 2011](#)).

Proposition 2.3.2 (Bounded perturbations for T_2). *Suppose that π satisfies a T_2 inequality. If $\mu \in \mathcal{P}_2(\mathbb{R}^d)$ is of the form $\frac{d\mu}{d\pi} \propto \exp B$, where $\text{osc } B \leq b$, then μ also satisfies a T_2 inequality with $C_{\text{T}_2}(\mu) \leq 8C_{\text{T}_2}(\pi) \exp b$.*

2.3.2 Lipschitz Mapping

Another simple but useful condition which enables us to transfer functional inequalities from π to μ is the existence of a Lipschitz mapping which pushes forward π to μ .

Proposition 2.3.3 (Lipschitz mapping). *Suppose that $\pi \in \mathcal{P}(\mathbb{R}^d)$ satisfies either a Poincaré or a log-Sobolev inequality, and that there exists an L -Lipschitz map $T : \mathbb{R}^d \rightarrow \mathbb{R}^d$ such that $\mu = T_{\#}\pi$. Then, μ also satisfies the corresponding functional inequality with constant*

$$C_{\text{PI}}(\mu) \leq L^2 C_{\text{PI}}(\pi) \quad \text{or} \quad C_{\text{LSI}}(\mu) \leq L^2 C_{\text{LSI}}(\pi)$$

respectively.

Proof Assume for simplicity that T is continuously differentiable, so that $\|\nabla T\|_{\text{op}} \leq L$. Then, for $f : \mathbb{R}^d \rightarrow \mathbb{R}$, by applying the Poincaré inequality for π ,

$$\begin{aligned} \text{var}_{\mu} f &= \text{var}_{\pi}(f \circ T) \leq C_{\text{PI}}(\pi) \mathbb{E}_{\pi}[\|\nabla(f \circ T)\|^2] \leq C_{\text{PI}}(\pi) \mathbb{E}_{\pi}[\|\nabla T\|_{\text{op}}^2 \|\nabla f \circ T\|^2] \\ &\leq C_{\text{PI}}(\pi) L^2 \mathbb{E}_{\pi}[\|\nabla f \circ T\|^2] = C_{\text{PI}}(\pi) L^2 \mathbb{E}_{\mu}[\|\nabla f\|^2]. \end{aligned}$$

The proof for the log-Sobolev inequality is similar. $\square$

This result becomes particularly powerful when combined with **Caffarelli's contraction theorem**, which states that the optimal transport map from the standard Gaussian to an α -strongly log-concave measure is $\alpha^{-1/2}$ -Lipschitz. As it is often easier to prove functional inequalities for the standard Gaussian, this principle then quickly implies Poincaré and log-Sobolev inequalities (as well as many other functional inequalities) for strongly log-concave measures. We will return to this in Section 3.5.

2.3.3 Lipschitz Perturbations

Whereas Proposition 2.3.1 implies that the LSI for strongly log-concave measures is robust against bounded perturbations, the following result shows that it is also robust against Lipschitz perturbations.

Proposition 2.3.4 (Lipschitz perturbations). *Let $\pi \propto \exp(-V - W)$, where V is α -strongly convex ($\alpha > 0$) and W is L -Lipschitz. Then,*

$$C_{\text{LSI}}(\pi) \leq \inf_{0 < u < \alpha} \frac{1}{u} \exp\left(\frac{L^2}{2(\alpha - u)}\right) \leq \frac{2}{\alpha} \exp\frac{L^2}{\alpha}.$$

Proof Let $V_u := V - (u/2) \|\cdot\|^2$ and note that V_u is $(\alpha - u)$ -strongly convex. Let $F_u := V_u + W$, and let $\text{conv } F_u$ denote the convex envelope of F_u :

$$\text{conv } F_u(x) := \inf \left\{ \sum_{i=1}^n \lambda_i F_u(x_i) : n \in \mathbb{N}, \lambda \in \mathbb{R}_+^n, \sum_{i=1}^n \lambda_i = 1, \sum_{i=1}^n \lambda_i x_i = x \right\}.$$

By definition, we have $\text{conv } F_u \leq F_u$. Also, for $\sum_{i=1}^n \lambda_i x_i = x$,

$$\begin{aligned} F_u(x) - \sum_{i=1}^n \lambda_i F_u(x_i) &= V_u(x) - \sum_{i=1}^n \lambda_i V_u(x_i) + \sum_{i=1}^n \lambda_i (W(x) - W(x_i)) \\ &\leq -\frac{\alpha - u}{2} \sum_{i=1}^n \lambda_i \|x_i - x\|^2 + L \sum_{i=1}^n \lambda_i \|x_i - x\| \leq \frac{L^2}{2(\alpha - u)}. \end{aligned}$$

Taking the infimum over $\lambda_1, \dots, \lambda_n$ and $x_1, \dots, x_n$ yields $F_u \leq \text{conv } F_u + L^2/2(\alpha - u)$. Thus, we have the decomposition

$$V + W = \underbrace{\text{conv } F_u + \frac{u}{2} \|\cdot\|^2}_{u\text{-strongly convex}} + \underbrace{V_u + W - \text{conv } F_u}_{\text{osc}(\cdot) \leq L^2/(2(\alpha - u))}.$$

The result follows from the Holley–Stroock principle ([Proposition 2.3.1](#)). □

2.3.4 Convolution

Next, we show that if $\pi = \pi_1 * \pi_2$ is a convolution of two measures, where both π_1 and π_2 satisfy a functional inequality, then so does π . This is a consequence of the subadditivity of variance and entropy. We begin with a variational principle for the entropy.

Lemma 2.3.5 (Variational principle for entropy). *For $f : \mathbb{R}^d \rightarrow \mathbb{R}_{>0}$,*

$$\text{ent}_\pi f = \sup \{ \mathbb{E}_\pi [fg] \mid g : \mathbb{R}^d \rightarrow \mathbb{R} \text{ such that } \mathbb{E}_\pi \exp g \leq 1 \}.$$

Proof We may assume that $\mathbb{E}_\pi \exp g = 1$, and define μ via $\frac{d\mu}{d\pi} := \exp g$. Then,

$$\text{ent}_\pi f = \mathbb{E}_\pi \left[f \log \frac{f}{\mathbb{E}_\pi f} \right] = \underbrace{\mathbb{E}_\mu \left[f \exp(-g) \log \frac{f \exp(-g)}{\mathbb{E}_\mu [f \exp(-g)]} \right]}_{=\text{ent}_\mu(f \exp(-g)) \geq 0} + \mathbb{E}_\pi [fg].$$

Equality holds if $g = \log(f/\mathbb{E}_\pi f)$. □

Remark 2.3.6. The variational principle above is essentially a reformulation of the Donsker–Varadhan variational principle ([Theorem 1.5.7](#)).

Lemma 2.3.7 (Subadditivity of variance and entropy). *If $X_1, \dots, X_n$ are independent random variables, then*

$$\begin{aligned} \text{var } f(X_1, \dots, X_n) &\leq \mathbb{E} \sum_{i=1}^n \text{var}(f(X_1, \dots, X_n) \mid X_{-i}), \\ \text{ent } f(X_1, \dots, X_n) &\leq \mathbb{E} \sum_{i=1}^n \text{ent}(f(X_1, \dots, X_n) \mid X_{-i}), \quad f \geq 0. \end{aligned}$$

Here, $\text{var}(\cdot \mid X_{-i})$ and $\text{ent}(\cdot \mid X_{-i})$ denote the conditional variance and entropy respectively when all variables except X_i are held fixed, i.e.,

$$\begin{aligned} \text{var}(f(X_1, \dots, X_n) \mid X_{-i} = x_{-i}) &= \text{var } f(x_1, \dots, x_{i-1}, X_i, x_{i+1}, \dots, x_n), \\ \text{ent}(f(X_1, \dots, X_n) \mid X_{-i} = x_{-i}) &= \text{ent } f(x_1, \dots, x_{i-1}, X_i, x_{i+1}, \dots, x_n). \end{aligned}$$

Proof The subadditivity of the variance was established in [Exercise 1.1](#), so we turn to entropy. Let $Z := f(X_1, \dots, X_n)$ for a positive function f , and

$$\Delta_i := \log \mathbb{E}[Z \mid X_1, \dots, X_i] - \log \mathbb{E}[Z \mid X_1, \dots, X_{i-1}],$$

so that

$$\text{ent } Z = \mathbb{E}[Z (\log Z - \log \mathbb{E} Z)] = \mathbb{E}\left[Z \sum_{i=1}^n \Delta_i\right].$$

Since

$$\mathbb{E}[\exp \Delta_i \mid X_{-i}] = \frac{\mathbb{E}[\mathbb{E}[Z \mid X_1, \dots, X_i] \mid X_{-i}]}{\mathbb{E}[Z \mid X_1, \dots, X_{i-1}]} = 1,$$

the variational principle yields

$$\mathbb{E}[Z \Delta_i] = \mathbb{E} \mathbb{E}[Z \Delta_i \mid X_{-i}] \leq \mathbb{E} \text{ent}(Z \mid X_{-i}).$$

□

Proposition 2.3.8 (Functional inequalities and convolutions). *Suppose that $\pi = \pi_1 * \pi_2 \in \mathcal{P}(\mathbb{R}^d)$, where π_1 and π_2 both satisfy either a Poincaré or a log-Sobolev inequality. Then, π also satisfies the corresponding functional inequality with constant*

$$C_{\text{PI}}(\pi) \leq C_{\text{PI}}(\pi_1) + C_{\text{PI}}(\pi_2) \quad \text{or} \quad C_{\text{LSI}}(\pi) \leq C_{\text{LSI}}(\pi_1) + C_{\text{LSI}}(\pi_2)$$

respectively.

Proof Let $X \sim \pi_1$ and $Y \sim \pi_2$ be independent, and let $f : \mathbb{R}^d \rightarrow \mathbb{R}$. Using the subadditivity of the variance ([Lemma 2.3.7](#)),

$$\begin{aligned} \text{var}_\pi f &= \text{var } f(X + Y) \leq \mathbb{E} \text{var}(f(X + Y) \mid Y) + \mathbb{E} \text{var}(f(X + Y) \mid X) \\ &\leq \{C_{\text{PI}}(\pi_1) + C_{\text{PI}}(\pi_2)\} \mathbb{E}[\|\nabla f(X + Y)\|^2], \end{aligned}$$

and a similar argument holds for the entropy. □

2.3.5 Mixtures

Suppose that π is a mixture, $\pi = \mu P := \int P_x \mu(dx)$, where $\mu \in \mathcal{P}(\mathcal{X})$ is the *mixing measure* and $(P_x)_{x \in \mathcal{X}}$ is a family of probability measures on $\mathbb{R}^d$ indexed by $\mathcal{X}$ (in other words, a Markov kernel). For example, when $\mathcal{X} = [k]$, then μP is a mixture of k distributions $P_1, \dots, P_k$ with mixing weights given by μ . When $\mathcal{X} = \mathbb{R}^d$ and P_x is the translation of a fixed probability measure $\nu \in \mathcal{P}(\mathbb{R}^d)$ by x , then $\mu P = \mu * \nu$ is the convolution of μ and ν .

Under general conditions on the mixture, it turns out that if each P_x satisfies a functional inequality, then so does the mixture μP . The simplest demonstration of this idea is for the Poincaré inequality. Although the arguments in this section apply more generally to mixtures μP on arbitrary state spaces, we focus on the $\mathbb{R}^d$ case for simplicity.

Proposition 2.3.9 (PI for mixtures, [Bardet et al. \(2018\)](#)). *Let μP be a mixture and assume that each P_x satisfies a Poincaré inequality with constant $C_{\text{PI}}(P)$. Also, assume that*

$$C_{\chi^2} := \sup_{x, x' \in \text{supp}(\mu)} \chi^2(P_x \parallel P_{x'}) < \infty. \quad (2.3.10)$$

Then, μP satisfies a Poincaré inequality with constant

$$C_{\text{PI}}(\mu P) \leq \left(1 + \frac{C_{\chi^2}}{4}\right) C_{\text{PI}}(P).$$

Proof Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$, and let $X \sim \mu$. By the law of total variance,

$$\text{var}_{\mu P} f = \mathbb{E} \text{var}_{P_X} f + \text{var} \mathbb{E}_{P_X} f.$$

The first term is easy to control, because we can apply the Poincaré inequality for P_X inside the expectation, so the main difficulty lies in the second term. Note that for any $x, x' \in \text{supp} \mu$,

$$|\mathbb{E}_{P_{x'}} f - \mathbb{E}_{P_x} f|^2 = \left| \int f \left(\frac{dP_{x'}}{dP_x} - 1 \right) dP_x \right|^2 \leq (\text{var}_{P_x} f) \chi^2(P_{x'} \parallel P_x) \leq C_{\chi^2} \text{var}_{P_x} f.$$

Let $e_- := \inf_{x \in \text{supp} \mu} \mathbb{E}_{P_x} f$ and $e_+ := \sup_{x \in \text{supp} \mu} \mathbb{E}_{P_x} f$. This implies

$$\text{var} \mathbb{E}_{P_X} f \leq \mathbb{E} \left[\left| \mathbb{E}_{P_X} f - \frac{e_- + e_+}{2} \right|^2 \right] \leq \frac{1}{4} \mathbb{E} \max\{(\mathbb{E}_{P_X} f - e_-)^2, (\mathbb{E}_{P_X} f - e_+)^2\} \leq \frac{C_{\chi^2}}{4} \mathbb{E} \text{var}_{P_X} f.$$

Hence,

$$\text{var}_{\mu P} f \leq \left(1 + \frac{C_{\chi^2}}{4}\right) \mathbb{E} \text{var}_{P_X} f \leq \left(1 + \frac{C_{\chi^2}}{4}\right) C_{\text{PI}}(P) \underbrace{\mathbb{E} \mathbb{E}_{P_X} \Gamma(f)}_{= \mathbb{E}_{\mu P} \Gamma(f)}. \quad \square$$

Our aim is to extend this idea to the log-Sobolev inequality, which will require a few preliminaries. Rather than aiming to directly prove a log-Sobolev inequality, we will instead prove a **defective log-Sobolev inequality**: for all $f : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\text{ent}_{\pi}(f^2) \leq 2C \mathbb{E}_{\pi} \Gamma(f) + D \mathbb{E}_{\pi}[f^2]. \quad (2.3.11)$$

Although the defective LSI involves an extra term on the right-hand side of the inequality, the extra term can be removed via a Poincaré inequality. The following standard result is proven via a lemma of O. Rothaus ([Rothaus, 1985](#)), see [Bakry et al. \(2014, Proposition 5.1.3\)](#).

Lemma 2.3.12 (Tightening a defective LSI). *Suppose that π satisfies the defective log-Sobolev inequality (2.3.11), together with a Poincaré inequality. Then, π satisfies a log-Sobolev inequality with constant*

$$C_{\text{LSI}} \leq C + \left(1 + \frac{D}{2}\right) C_{\text{PI}}.$$

We also need one change of measure lemma.

Lemma 2.3.13 (Chen et al. (2021a)). *Suppose that μ, ν are probability measures and f is a positive function. Then,*

$$\mathbb{E}_\mu(f) \log \frac{\mathbb{E}_\mu(f)}{\mathbb{E}_\nu(f)} \leq \text{ent}_\mu(f) + \mathbb{E}_\mu(f) \log(1 + \chi^2(\mu \parallel \nu)).$$

Proof We can assume $\mu \ll \nu$, or else the statement is trivial. By rescaling, we may assume $\mathbb{E}_\mu f = 1$. Recall the Donsker–Varadhan variational principle (Theorem 1.5.7), which states

$$\text{KL}(\eta \parallel \eta') = \sup \{ \mathbb{E}_\eta g - \log \mathbb{E}_{\eta'} \exp g \mid g : \mathcal{X} \rightarrow \mathbb{R} \text{ is bounded and measurable} \}.$$

If we take $\eta = f\mu$, $\eta' = \nu$, and $g = \log(f/\mathbb{E}_\nu f)$, then

$$\mathbb{E}_\mu \left[f \log \frac{f}{\mathbb{E}_\nu f} \right] = \mathbb{E}_\eta \log \frac{f}{\mathbb{E}_\nu f} \leq \text{KL}(\eta \parallel \nu) + \underbrace{\log \mathbb{E}_\nu \frac{f}{\mathbb{E}_\nu f}}_{=0} = \mathbb{E}_\mu \left[f \log \left(f \frac{d\mu}{d\nu} \right) \right].$$

By subtracting $\mathbb{E}_\mu(f \log f)$ from both sides, we obtain

$$\log \frac{1}{\mathbb{E}_\nu f} = \mathbb{E}_\mu \left[f \log \frac{1}{\mathbb{E}_\nu f} \right] \leq \mathbb{E}_\mu \left[f \log \frac{d\mu}{d\nu} \right].$$

Next, applying the Donsker–Varadhan variational principle a second time with $\eta = f\mu$, $\eta' = \mu$, and $g = \log \frac{d\mu}{d\nu}$ yields

$$\mathbb{E}_\mu \left[f \log \frac{d\mu}{d\nu} \right] = \mathbb{E}_\eta \log \frac{d\mu}{d\nu} \leq \text{KL}(\eta \parallel \mu) + \log \mathbb{E}_\mu \frac{d\mu}{d\nu} = \text{ent}_\mu f + \log(1 + \chi^2(\mu \parallel \nu)),$$

which is what we wanted to show. $\square$

We can now prove the log-Sobolev inequality for mixtures.

Proposition 2.3.14 (LSI for mixtures, Chen et al. (2021a)). *Let μP be a mixture and assume that each P_x satisfies a log-Sobolev inequality with constant $C_{\text{LSI}}(P)$. Also, assume that*

$$C_{\chi^2} := \sup_{x, x' \in \text{supp}(\mu)} \chi^2(P_x \parallel P_{x'}) < \infty.$$

Then, μP satisfies a log-Sobolev inequality with constant

$$C_{\text{LSI}}(\mu P) \leq C_{\text{LSI}}(P) \left\{ 2 + \left(1 + \frac{C_{\chi^2}}{4}\right) \left(1 + \frac{1}{2} \log(1 + C_{\chi^2})\right) \right\}.$$

Proof We begin, as in the proof of [Proposition 2.3.9](#), with a decomposition of the entropy:

$$\text{ent}_{\mu P}(f^2) = \mathbb{E} \text{ent}_{P_X}(f^2) + \text{ent} \mathbb{E}_{P_X}(f^2).$$

As before, it is the second term that is difficult to control.

Applying [Lemma 2.3.13](#) and joint convexity of χ^2 ,

$$\begin{aligned} \text{ent} \mathbb{E}_{P_X}(f^2) &= \mathbb{E} \left[\mathbb{E}_{P_X}(f^2) \log \frac{\mathbb{E}_{P_X}(f^2)}{\mathbb{E}_{\mu P}(f^2)} \right] \leq \mathbb{E} \left[\text{ent}_{P_X}(f^2) + \mathbb{E}_{P_X}(f^2) \log(1 + \chi^2(P_X \parallel \mu P)) \right] \\ &\leq \mathbb{E} \text{ent}_{P_X}(f^2) + \log(1 + C_{\chi^2}) \mathbb{E}_{\mu P}(f^2). \end{aligned}$$

Hence,

$$\begin{aligned} \text{ent}_{\mu P}(f^2) &\leq 2 \mathbb{E} \text{ent}_{P_X}(f^2) + \log(1 + C_{\chi^2}) \mathbb{E}_{\mu P}(f^2) \\ &\leq 4C_{\text{LSI}}(P) \mathbb{E}_{\mu P} \Gamma(f) + \log(1 + C_{\chi^2}) \mathbb{E}_{\mu P}(f^2), \end{aligned}$$

where we have applied the log-Sobolev inequality for P_X . This is a defective log-Sobolev inequality for μP ; by applying the Poincaré inequality from [Proposition 2.3.9](#) and tightening the inequality via [Lemma 2.3.12](#), we conclude the proof. $\square$

Example 2.3.15 (LSI for Gaussian mixtures). Suppose that μ is supported on a ball $\mathbb{B}(0, R)$, and that for each $x \in \mathbb{R}^d$, $P_x = \text{normal}(x, \sigma^2 I_d)$. Then, μP is the convolution $\mu * \text{normal}(0, \sigma^2 I_d)$. Since $C_{\text{LSI}}(P) = \sigma^2$ and

$$\chi^2(P_x \parallel P_{x'}) = \exp \frac{\|x - x'\|^2}{\sigma^2} - 1 \leq \exp \frac{4R^2}{\sigma^2} - 1$$

for $x, x' \in \mathbb{B}(0, R)$, we deduce that μP satisfies a log-Sobolev inequality with constant

$$C_{\text{LSI}}(\mu P) \leq 3(\sigma^2 + 2R^2) \exp \frac{4R^2}{\sigma^2}.$$

Hence, *Gaussian convolutions of measures with bounded support satisfy a log-Sobolev inequality*. The exponential dependence on R^2/σ^2 is unavoidable in general. Another approach to the log-Sobolev inequality for Gaussian mixtures is given in [Exercise 3.9](#).

We extend the results of this section in [Exercise 2.14](#).

2.3.6 Tensorization

A key feature of these functional inequalities which makes them crucial for the study of high-dimensional, or even infinite-dimensional, phenomena is that they often hold with dimension-free constants, as demonstrated in the next result.

Theorem 2.3.16 (Tensorization). *Suppose that $\pi_1, \dots, \pi_N \in \mathcal{P}(\mathbb{R}^d)$ satisfy either a Poincaré inequality or a log-Sobolev inequality. Then, for any $N \in \mathbb{N}^+$, the product measure $\pi := \bigotimes_{i=1}^N \pi_i$ also satisfies the corresponding functional inequality with constant*

$$C_{\text{PI}}(\pi) = \max_{i \in [N]} C_{\text{PI}}(\pi_i) \quad \text{or} \quad C_{\text{LSI}}(\pi) = \max_{i \in [N]} C_{\text{LSI}}(\pi_i)$$

respectively.

Proof The proof is a straightforward consequence of subadditivity (Lemma 2.3.7). Indeed, if $f : \mathbb{R}^{Nd} \rightarrow \mathbb{R}$ and if $X_i \sim \pi_i$ are independent for $i \in [N]$,

$$\begin{aligned} \text{var}_\pi f &= \text{var} f(X_1, \dots, X_N) \leq \mathbb{E} \sum_{i=1}^N \text{var}(f(X_1, \dots, X_N) \mid X_{-i}) \\ &\leq \max_{i \in [N]} C_{\text{PI}}(\pi_i) \mathbb{E} \sum_{i=1}^N \mathbb{E}[\|\nabla_i f(X_1, \dots, X_N)\|^2 \mid X_{-i}] \\ &= \max_{i \in [N]} C_{\text{PI}}(\pi_i) \mathbb{E}_\pi[\|\nabla f\|^2]. \end{aligned}$$

The proof is the same for the log-Sobolev inequality. $\square$

There is also a tensorization principle for transport inequalities, which however requires some additional work to prove. We formulate a general result which applies to many different transport inequalities (not just the T_1 and T_2 inequalities). See the supplementary material for a proof.

Theorem 2.3.17 (Marton's tensorization, Marton (1996)). *Let $\mathcal{X}_1, \dots, \mathcal{X}_N$ be Polish spaces equipped with probability measures $\pi_1, \dots, \pi_N$ respectively. Let $\mathcal{X} := \mathcal{X}_1 \times \dots \times \mathcal{X}_N$ be equipped with the product measure $\pi := \pi_1 \otimes \dots \otimes \pi_N$.*

Let $\varphi : [0, \infty) \rightarrow [0, \infty)$ be convex and for $i \in [N]$, let $c_i : \mathcal{X}_i \times \mathcal{X}_i \rightarrow [0, \infty)$ be a lower semicontinuous cost function. Suppose that

$$\inf_{\gamma_i \in \mathcal{C}(\pi_i, \nu_i)} \varphi\left(\int c_i d\gamma_i\right) \leq 2\sigma^2 \text{KL}(\nu_i \parallel \pi_i), \quad \forall \nu_i \in \mathcal{P}(\mathcal{X}_i), \forall i \in [N].$$

Then, it holds that

$$\inf_{\gamma \in \mathcal{C}(\pi, \nu)} \sum_{i=1}^N \varphi\left(\int c_i(x_i, y_i) \gamma(dx_{1:N}, dy_{1:N})\right) \leq 2\sigma^2 \text{KL}(\nu \parallel \pi), \quad \forall \nu \in \mathcal{P}(\mathcal{X}).$$

We can apply Marton's tensorization (Theorem 2.3.17) with the convex function $\varphi(x) := x$ and cost functions $c_i := d_i^2$, with d_i a lower semicontinuous metric on $\mathcal{X}_i$. This immediately yields the following corollary.

Corollary 2.3.18 (Tensorization of T_2). *Suppose that for each $i \in [N]$, $\pi_i \in \mathcal{P}(\mathcal{X}_i)$ satisfies a T_2 inequality with parameter σ^2 with respect to the metric d_i . Then, the product measure $\pi_1 \otimes \dots \otimes \pi_N$ satisfies a T_2 inequality with the same parameter σ^2 with respect to the metric $d(x_{1:N}, y_{1:N})^2 := \sum_{i=1}^N d(x_i, y_i)^2$ on $\mathcal{X}_1 \times \dots \times \mathcal{X}_N$.*

2.4 Concentration of Measure and Isoperimetry

Arguably, the most prominent application of functional inequalities has been to the concentration of measure phenomenon, which is an indispensable tool in high-dimensional probability and statistics. Since it is not the focus of this book, we only mention a few key results, focusing on the connection with *isoperimetric inequalities* due to their relevance to Chapter 7. See, however, the supplementary material to this book for more details.

Since many of the arguments hold on a general Polish space (that is, a complete separable metric space) $(\mathcal{X}, d)$ equipped with a probability measure π , we will work in this setting unless explicitly stated otherwise.

2.4.1 Blow-Up of Sets and Concentration of Lipschitz Functions

Loosely speaking, the concentration of measure phenomenon holds when a huge fraction of the mass of π is concentrated on a relatively small set. Another way of capturing this idea is to assert that whenever a set has a non-trivial amount of mass under π , then expanding the set slightly causes it to capture *almost all* of the mass of π . The following definitions formalize this idea.

Definition 2.4.1. For a Borel subset $A \subseteq \mathcal{X}$ and $\varepsilon > 0$, we let A^ε denote the ε -**blow-up** of A , defined by

$$A^\varepsilon := \{x \in \mathcal{X} \mid d(x, A) < \varepsilon\}.$$

The **concentration function** $\alpha_\pi : \mathbb{R}_+ \rightarrow [0, 1]$ is defined via

$$\alpha_\pi(\varepsilon) = \sup\{\pi((A^\varepsilon)^c) \mid A \subseteq \mathcal{X} \text{ is a Borel subset with } \pi(A) \geq \frac{1}{2}\}.$$

Typically, we have $\alpha_\pi(\varepsilon) \leq C_0 \exp(-\varepsilon/C_1)$ or $\alpha_\pi(\varepsilon) \leq C_0 \exp(-\varepsilon^2/C_1)$ for some constants $C_0, C_1 > 0$; hence, as we increase ε , the blow-up A^ε captures an often substantially larger fraction of the mass of π . We first develop an equivalence between the formulation of concentration of measure via blow-up of sets, and with another involving concentration of Lipschitz functions. Although the former has a more striking geometric interpretation, the latter is often more useful for applications.

Given a real-valued random variable X , we abuse notation and let $\text{med } X$ denote any median of X , that is, any number m such that $\mathbb{P}\{X \leq m\} \wedge \mathbb{P}\{X \geq m\} \geq \frac{1}{2}$.

Theorem 2.4.2 (Blow-up and Lipschitz functions). *Suppose that $(\mathcal{X}, d, \pi)$ has concentration function α_π . Then, for any 1-Lipschitz function $f : \mathcal{X} \rightarrow \mathbb{R}$ and $\varepsilon \geq 0$,*

$$\pi\{f \geq \text{med } f + \varepsilon\} \leq \alpha_\pi(\varepsilon).$$

Conversely, suppose that for all 1-Lipschitz functions $f : \mathcal{X} \rightarrow \mathbb{R}$, it holds that

$$\pi\{f \geq \text{med } f + \varepsilon\} \leq \beta(\varepsilon).$$

Then, the concentration function α_π of $(\mathcal{X}, d, \pi)$ satisfies $\alpha_\pi \leq \beta$.

Proof ($\implies$) Consider the set $A := \{f \leq \text{med } f\}$. We claim that $A^\varepsilon \subseteq \{f - \text{med } f < \varepsilon\}$. To prove this, let $x \in A^\varepsilon$. By definition, there exists $y \in A$ such that $d(x, y) < \varepsilon$, so $f(x) - \text{med } f =$

$f(y) - \text{med } f + f(x) - f(y) \leq d(x, y) < \varepsilon$. Hence,

$$\pi\{f - \text{med } f < \varepsilon\} \geq \pi(A^\varepsilon) \geq 1 - \alpha_\pi(\varepsilon).$$

($\Leftarrow$) The function $f := d(\cdot, A)$ is 1-Lipschitz, and if $\pi(A) \geq \frac{1}{2}$ then 0 is a median of f . Thus, it holds that

$$\pi(A^\varepsilon) = \pi\{f - \text{med } f < \varepsilon\} \geq 1 - \beta(\varepsilon).$$

□

More broadly, it is a general principle that many statements about sets have an equivalent reformulation in terms of functions.

Whether by bounding the concentration function, or by establishing concentration inequalities for Lipschitz functions directly, by now there is a wealth of techniques for proving results of the following type. We remark that while the following theorem asserts concentration around the mean, whereas [Theorem 2.4.2](#) considers concentration around the median, the difference is merely superficial since the mean and median are essentially equivalent (see, e.g., [Milman, 2009](#)).

Theorem 2.4.3 (Functional inequalities and concentration). *Let $\pi \in \mathcal{P}(\mathbb{R}^d)$, and let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be a 1-Lipschitz function.*

1 *If π satisfies a Poincaré inequality with constant C_{PI} , then for all $t \geq 0$,*

$$\pi\{f - \mathbb{E}_\pi f \geq t\} \leq 3 \exp\left(-\frac{2t}{\sqrt{C_{\text{PI}}}}\right).$$

2 *If π satisfies Talagrand's T_1 inequality with constant C_{T_1} , then for all $t \geq 0$,*

$$\pi\{f - \mathbb{E}_\pi f \geq t\} \leq \exp\left(-\frac{t^2}{2C_{T_1}}\right).$$

In particular, recalling that $C_{T_1} \leq C_{T_2} \leq C_{\text{LSI}}$, the second statement furnishes a concentration inequality under a log-Sobolev inequality. Thus, a Poincaré inequality implies subexponential concentration, and a log-Sobolev inequality implies sub-Gaussian concentration. Moreover, due to tensorization results (Section 2.3.6), the constants involved are dimension-free for many situations of interest. Consequently, [Theorem 2.4.3](#) is immensely useful for high-dimensional applications. Here is a basic illustration.

Example 2.4.4. Suppose that γ is the standard Gaussian measure on $\mathbb{R}^d$. From the Bakry–Émery theorem ([Theorem 1.2.30](#)), γ satisfies the log-Sobolev inequality with $C_{\text{LSI}} = 1$. For $Z \sim \gamma$, since $\mathbb{E}[\|Z\|^2] = d$, the Poincaré inequality applied to the norm $\|\cdot\|$ shows that $\text{var } \|Z\| \leq 1$, i.e., $\sqrt{d-1} \leq \mathbb{E}\|Z\| \leq \sqrt{d}$.

The concentration result above now shows that the standard Gaussian “lives” on a thin spherical shell of radius $\sqrt{d}$ and width $O(1)$.

2.4.2 Classical Isoperimetry Results

We now study finer questions about the concentration function. More generally, for $p \in (0, 1)$ and $\varepsilon > 0$, we introduce the quantity

$$\omega_\pi(p, \varepsilon) := \inf\{\pi(A^\varepsilon) \mid A \text{ is a Borel set with } \pi(A) = p\},$$

and we can ask about the deviation of $\omega_\pi(p, \varepsilon)$ from p when ε is small. In some special cases, we can even determine the function ω_π exactly. The study of this question will bring us to the classical geometric problem of isoperimetry. In its simplest guise, it asks: among all plane curves which enclose an area of a prescribed area, which ones have the least perimeter? Unsurprisingly, among regular curves, it is well-known that circles provide the answer to this question. As we shall see, the isoperimetric question, once generalized to abstract spaces, contains a wealth of information about concentration phenomena.

The connection between concentration and isoperimetry began with the work of Lévy, who found the isoperimetric inequality on the sphere.

Theorem 2.4.5 (Spherical isoperimetry). *Let σ_d denote the uniform measure on the d -dimensional unit sphere $\mathbb{S}^d$, and let $A \subseteq \mathbb{S}^d$ be a Borel subset with $\sigma_d(A) \notin \{0, 1\}$. Let C be a spherical cap with the same measure as A . Then, for all $\varepsilon > 0$,*

$$\sigma_d(A^\varepsilon) \geq \sigma_d(C^\varepsilon).$$

We give a few reminders about spherical geometry. We equip $\mathbb{S}^d$ with its geodesic metric $\mathbf{d}$, so that the distance between two points $x, y \in \mathbb{S}^d$ is equal to the angle between x and y . A spherical cap is a geodesic ball, that is, it is a set of the form $B(x_0, r)$ for some $x_0 \in \mathbb{S}^d$ and $r > 0$, where the balls are defined w.r.t. $\mathbf{d}$.

Note that [Theorem 2.4.5](#) identifies the exact function ω_{σ_d} . To see how an isoperimetric result naturally leads to a concentration result, suppose that A has measure $\frac{1}{2}$. Then, the corresponding spherical cap C can be taken to be half of the sphere, $C = B(x_0, \frac{\pi}{2})$, and so

$$\sigma_d(A^\varepsilon) \geq \sigma_d(C^\varepsilon) = \sigma_d\left(B\left(x_0, \frac{\pi}{2} + \varepsilon\right)\right).$$

To obtain an upper bound on the concentration function α_{σ_d} , it therefore suffices to lower bound the volume of the spherical cap. It leads to the following result, which we leave as [Exercise 2.18](#).

Theorem 2.4.6 (Concentration on the sphere). *Let σ_d be the uniform measure on the unit sphere $\mathbb{S}^d$ in dimension $d \geq 2$, equipped with the geodesic distance $\mathbf{d}$. Then,*

$$\alpha_{\sigma_d}(\varepsilon) \leq \exp\left(-\frac{(d-1)\varepsilon^2}{2}\right), \quad \text{for all } \varepsilon > 0.$$

There is also an isoperimetric result for the standard Gaussian measure γ_d on $\mathbb{R}^d$. In this case, the optimal sets are given by half-spaces, i.e., sets of the form

$$H_{x_0, t} := \{x \in \mathbb{R}^d \mid \langle x, x_0 \rangle \leq t\}.$$

Theorem 2.4.7 (Gaussian isoperimetry). *Let γ_d denote the standard Gaussian measure on $\mathbb{R}^d$, and let $A \subseteq \mathbb{R}^d$ be a Borel subset with $\gamma_d(A) \notin \{0, 1\}$. Let $H_{x_0, t}$ be a half-space with the same measure as A . Then, for all $\varepsilon > 0$,*

$$\gamma_d(A^\varepsilon) \geq \gamma_d(H_{x_0, t}^\varepsilon).$$

We can write this result more explicitly as follows. By rotational invariance of the Gaussian, we can take x_0 to be any unit vector e , in which case the measure of $H_{e, t}$ is $\gamma_d(H_{e, t}) = \Phi(t)$, where Φ is the Gaussian CDF. Since $H_{e, t}^\varepsilon = H_{e, t+\varepsilon}$, then

$$\gamma_d(A^\varepsilon) \geq \Phi(\Phi^{-1}(\gamma_d(A)) + \varepsilon). \quad (2.4.8)$$

In particular, if $\gamma_d(A) = \frac{1}{2}$, then $\Phi^{-1}(\frac{1}{2}) = 0$, so

$$\alpha_{\gamma_d}(\varepsilon) \leq \Phi(-\varepsilon) \leq \frac{1}{2} \exp\left(-\frac{\varepsilon^2}{2}\right).$$

We now pause to give a remark on proofs. Since these isoperimetric inequalities require a detailed understanding of the measure (including the optimal sets in the inequality), they are considerably more difficult to prove than the other results we have seen so far (e.g., a log-Sobolev inequality). In particular, usually they can only be established for measures which are simple in some regard, e.g., they enjoy many symmetries. Hence, we will not prove them here.

It is often convenient to pass to a differential form of the isoperimetric inequality, which is obtained by sending $\varepsilon \searrow 0$. This is formalized as follows.

Definition 2.4.9. Given a non-empty Borel set A and a measure π on a Polish space $(\mathcal{X}, d)$, the **Minkowski content** of A under π is

$$\pi^+(A) := \liminf_{\varepsilon \searrow 0} \frac{\pi(A^\varepsilon) - \pi(A)}{\varepsilon}.$$

Definition 2.4.10. For a measure π on a Polish space $(\mathcal{X}, d)$, the **isoperimetric profile** of π , denoted $\mathcal{I}_\pi : [0, 1] \rightarrow \mathbb{R}_+$, is the function

$$\mathcal{I}_\pi(p) := \inf\{\pi^+(A) \mid A \text{ is measurable with } \pi(A) = p\}.$$

For the standard Gaussian, a Taylor expansion of (2.4.8) yields

$$\gamma_d(A^\varepsilon) \geq \gamma_d(A) + \phi(\Phi^{-1}(\gamma_d(A))) \varepsilon + o(\varepsilon),$$

where $\phi = \Phi'$ is the Gaussian density. Hence, we can identify the isoperimetric profile of the standard Gaussian as

$$\mathcal{I}_{\gamma_d}(p) = \phi(\Phi^{-1}(p)). \quad (2.4.11)$$

Actually, (2.4.11) is *equivalent* to the Gaussian isoperimetric inequality (2.4.8). Here, (2.4.8) is called the *integral* form of the inequality, whereas (2.4.11) is called the *differential* form. The following theorem shows how to convert between the two forms.

Theorem 2.4.12 (Bobkov and Houdré (1997)). Let $I : (0, 1) \rightarrow \mathbb{R}_{>0}$. Define the increasing function F such that $F(0) = \frac{1}{2}$ and $f \circ F^{-1} = I$, where $f = F'$; equivalently, we can take $F^{-1}(p) = \int_{1/2}^p I(t)^{-1} dt$. Then, the following statements are equivalent.

1 For all $\varepsilon > 0$ and all Borel A with $\pi(A) \notin \{0, 1\}$,

$$\pi(A^\varepsilon) \geq F(F^{-1}(\pi(A)) + \varepsilon).$$

2 For all Borel A with $\pi(A) \notin \{0, 1\}$,

$$\pi^+(A) \geq I(\pi(A)).$$

Proof sketch Let $\bar{\omega}(p, \varepsilon) := F(F^{-1}(p) + \varepsilon)$. Then, one can show that $\bar{\omega}$ satisfies the semigroup property $\bar{\omega}(\bar{\omega}(p, \varepsilon), \varepsilon') = \bar{\omega}(p, \varepsilon + \varepsilon')$. Using this, one shows that to prove $\pi(A^\varepsilon) \geq \bar{\omega}(\pi(A), \varepsilon)$ it suffices to consider $\varepsilon \searrow 0$. A Taylor expansion yields

$$\begin{aligned} \pi(A^\varepsilon) &\geq \pi(A) + \pi^+(A) \varepsilon + o(\varepsilon), \\ \bar{\omega}(\pi(A), \varepsilon) &= \pi(A) + I(\pi(A)) \varepsilon + o(\varepsilon), \end{aligned}$$

from which we deduce that $\pi^+(A) \geq I(\pi(A))$ for all A if and only if $\pi(A^\varepsilon) \geq \bar{\omega}(\pi(A), \varepsilon)$ for all A and all $\varepsilon > 0$. $\square$

2.4.3 Cheeger Isoperimetry

We now consider a class of probability measures which is characterized by a lower bound on the isoperimetric profile.

Definition 2.4.13. A probability measure π satisfies a **Cheeger isoperimetric inequality** with constant $\text{Ch} > 0$ if for all Borel sets $A \subseteq \mathcal{X}$,

$$\pi^+(A) \geq \frac{1}{\text{Ch}} \pi(A) \pi(A^c). \quad (2.4.14)$$

For the two-sided exponential density $x \mapsto \mu(x) := \frac{1}{2} \exp(-|x|)$, the isoperimetric profile is known to be $\mathcal{I}_\mu(p) = \min(p, 1 - p)$. Hence, the Cheeger isoperimetric inequality roughly asserts that the isoperimetric properties of π are at least as good as those of μ .

The inequality in (2.4.14) is the differential form of the inequality. By applying Theorem 2.4.12, one shows that the inequality (2.4.14) implies, for any $\varepsilon \in [0, \text{Ch}]$,

$$\pi(A^\varepsilon) - \pi(A) \geq \frac{\varepsilon}{2 \text{Ch}} \pi(A) \pi(A^c). \quad (2.4.15)$$

For all $\varepsilon > 0$, Theorem 2.4.12 also implies that

$$\alpha_\pi(\varepsilon) \leq \exp\left(-\frac{\varepsilon}{\text{Ch}}\right),$$

so π enjoys at least subexponential concentration.

Such isoperimetric inequalities will play a key role when we study Metropolis-adjusted sampling algorithms in Chapter 7. For now, however, our goal is to establish an equivalence between the Cheeger isoperimetric inequality and a functional version of it.

To pass from a functional inequality to an inequality involving sets, we can usually apply the functional inequality to the indicator of a set. To go the other way around, we need to represent a function via its level sets, which is achieved via the coarea inequality.

Theorem 2.4.16 (Coarea inequality). *Let $f : \mathcal{X} \rightarrow \mathbb{R}$ be Lipschitz. Then,*

$$\int \|\nabla f\| d\pi \geq \int_{-\infty}^{\infty} \pi^+\{f > t\} dt.$$

Remark 2.4.17. On a general metric space $(\mathcal{X}, d)$, we define

$$\|\nabla f\|(x) := \limsup_{y \in \mathcal{X}, d(x,y) \searrow 0} \frac{|f(x) - f(y)|}{d(x,y)}.$$

In “nice” spaces, the coarea inequality is actually an equality, but we will not need this.

Proof sketch By an approximation argument we may assume that f is bounded, and by adding a constant to f we may suppose $f \geq 0$. Let $f_\varepsilon(x) := \sup_{d(x,\cdot) < \varepsilon} f$ and $A_t := \{f > t\}$. We can check that $A_t^\varepsilon = \{f_\varepsilon > t\}$, and for $g \geq 0$ we have the formula $\int g d\pi = \int_0^\infty \pi\{g > t\} dt$. By applying this to $g = f$ and $g = f_\varepsilon$,

$$\int \frac{f_\varepsilon - f}{\varepsilon} d\pi = \int_0^\infty \frac{\pi(A_t^\varepsilon) - \pi(A_t)}{\varepsilon} dt.$$

Now let $\varepsilon \searrow 0$ using Fatou’s lemma and dominated convergence. □

Theorem 2.4.18 (Functional form of Cheeger isoperimetry). *Let $\pi \in \mathcal{P}_1(\mathcal{X})$ and let $\text{Ch} > 0$. The following are equivalent.*

- 1 π satisfies a Cheeger isoperimetric inequality with constant Ch .
- 2 For all Lipschitz $f : \mathcal{X} \rightarrow \mathbb{R}$, it holds that

$$\mathbb{E}_\pi |f - \mathbb{E}_\pi f| \leq 2 \text{Ch} \mathbb{E}_\pi \|\nabla f\|. \quad (2.4.19)$$

Proof sketch (2) $\implies$ (1): Apply (2.4.19) to an approximation f of the indicator function $\mathbb{1}_A$, so that $\mathbb{E}_\pi |f - \mathbb{E}_\pi f| \approx 2 \pi(A) (1 - \pi(A))$ and $\mathbb{E}_\pi \|\nabla f\| \approx \pi^+(A)$.

(1) $\implies$ (2): By approximation, we can assume that f is non-negative. Let $A_t := \{f > t\}$. Applying the coarea inequality and the Cheeger isoperimetric inequality,

$$\begin{aligned} 2 \text{Ch} \mathbb{E}_\pi \|\nabla f\| &\geq 2 \text{Ch} \int_{-\infty}^{\infty} \pi^+\{f > t\} dt \\ &\geq 2 \int_{-\infty}^{\infty} \pi(A_t) \pi(A_t^c) dt = \int_{-\infty}^{\infty} \mathbb{E}_\pi |\mathbb{1}_{A_t} - \pi(A_t)| dt \\ &\geq \sup_{\|g\|_{L^\infty(\pi)} \leq 1} \int_{-\infty}^{\infty} \left(\int g \{\mathbb{1}_{A_t} - \pi(A_t)\} d\pi \right) dt \\ &= \sup_{\|g\|_{L^\infty(\pi)} \leq 1} \int_{-\infty}^{\infty} \left(\int \{g - \mathbb{E}_\pi g\} \mathbb{1}_{A_t} d\pi \right) dt = \sup_{\|g\|_{L^\infty(\pi)} \leq 1} \int \{g - \mathbb{E}_\pi g\} f d\pi \\ &= \mathbb{E}_\pi |f - \mathbb{E}_\pi f|. \end{aligned} \quad \square$$

2.4.4 L^p – L^q Poincaré Inequalities

In this section, we work on Euclidean space for simplicity.

The inequality (2.4.19) can be considered an “ L^1 variant” of the Poincaré inequality. More generally, we can define the following family of inequalities.

Definition 2.4.20 (L^p – L^q Poincaré inequality). For $p, q \in [1, \infty]$ with $q \geq p$, the L^p – L^q Poincaré inequality asserts that for all smooth $f : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\|f - \mathbb{E}_\pi f\|_{L^p(\pi)} \leq C_{p,q} \|\nabla f\|_{L^q(\pi)}.$$

In this new notation, the usual Poincaré inequality is an L^2 – L^2 Poincaré inequality with $C_{2,2} = \sqrt{C_{\text{PI}}}$, whereas the inequality (2.4.19) is an L^1 – L^1 Poincaré inequality.

These inequalities form a hierarchy via Hölder’s inequality.

Proposition 2.4.21 (Milman (2009)). Suppose $p, q, \hat{p}, \hat{q} \in [1, \infty]$ are such that $p \leq \hat{p}$ and $q \leq \hat{q}$, and $p^{-1} - q^{-1} = \hat{p}^{-1} - \hat{q}^{-1}$. Then,

$$C_{\hat{p}, \hat{q}} \lesssim \frac{\hat{p}}{p} C_{p,q}.$$

Proof Let f satisfy $\text{med}_\pi f = 0$, which we can arrange by adding a constant. Define the function $g := (\text{sgn } f) |f|^{\hat{p}/p}$, which still satisfies $\text{med}_\pi g = 0$. Throughout the proof, we use the fact that for any function h , the $L^p(\pi)$ norms of $h - \mathbb{E}_\pi h$ and $h - \text{med}_\pi h$ are equivalent up to a universal constant (Milman, 2009). Applying the L^p – L^q Poincaré inequality to g together with Hölder’s inequality,

$$\begin{aligned} \|f - \text{med}_\pi f\|_{L^{\hat{p}}(\pi)}^{\hat{p}/p} &= \|g - \text{med}_\pi g\|_{L^p(\pi)} \lesssim \|g - \mathbb{E}_\pi g\|_{L^p(\pi)} \leq C_{p,q} \|\nabla g\|_{L^q(\pi)} \\ &= \frac{\hat{p}}{p} C_{p,q} \| |f|^{\hat{p}/p-1} \nabla f \|_{L^q(\pi)} \\ &\leq \frac{\hat{p}}{p} C_{p,q} \|f - \text{med}_\pi f\|_{L^{\hat{p}}(\pi)}^{\hat{p}/p-1} \|\nabla f\|_{L^{\hat{q}}(\pi)}, \end{aligned}$$

where we leave it to the reader to check that the exponents work out correctly. If we rearrange this inequality, then

$$\|f - \mathbb{E}_\pi f\|_{L^{\hat{p}}(\pi)} \lesssim \|f - \text{med}_\pi f\|_{L^{\hat{p}}(\pi)} \lesssim \frac{\hat{p}}{p} C_{p,q} \|\nabla f\|_{L^{\hat{q}}(\pi)}.$$

□

Thus, we have the following implications: for any $p \in (2, \infty)$,

$$(L^1 - L^1) \implies (L^2 - L^2) \implies \dots \implies (L^p - L^p).$$

In particular, the first implication together with the equivalence in Theorem 2.4.18 shows that the Cheeger isoperimetric inequality implies the Poincaré inequality.

Also, given any L^p – L^q Poincaré inequality, by Jensen’s inequality we can trivially make it weaker by decreasing p or increasing q ; hence, every L^p – L^q Poincaré inequality implies an L^1 – L^∞ Poincaré inequality. On the other hand, for any $1 \leq p \leq q < \infty$, an L^1 – L^1 Poincaré implies an L^p – L^p Poincaré, which trivially implies an L^p – L^q Poincaré inequality. We conclude that *among these inequalities, the L^1 – L^1 inequality is the strongest and the L^1 – L^∞ inequality is the weakest.*

We now sketch the proof of a deep result by E. Milman, which states that for log-concave measures, the hierarchy can be reversed. The formal statement is as follows.

Theorem 2.4.22 (Reversing the hierarchy, [Milman \(2009\)](#)). *Let π be log-concave. Then,*

$$C_{1,1} \lesssim C_{1,\infty}.$$

As a consequence, suppose that π is α -strongly log-concave. By the Bakry–Émery theorem ([Theorem 1.2.30](#)), π satisfies a Poincaré inequality with $C_{2,2}^2 = C_{PI} \leq 1/\alpha$. By reversing the hierarchy, we see that this implies a Cheeger isoperimetric inequality.

Corollary 2.4.23. *If $\pi \in \mathcal{P}(\mathbb{R}^d)$ is α -strongly log-concave, then π satisfies a Cheeger isoperimetric inequality with constant $Ch \lesssim 1/\sqrt{\alpha}$.*

The proof of Milman’s theorem will require some preparations. The first fact that we need is a deep result in its own right. Typically it is proven with geometric measure theory by studying the isoperimetric problem, and we omit the proof; see [Milman \(2009\)](#).

Theorem 2.4.24. *If π is log-concave, then its isoperimetric profile $\mathcal{I}_\pi$ is concave.*

The isoperimetric profile satisfies $\mathcal{I}_\pi(0) = 0$, and it is symmetric around $\frac{1}{2}$, so it suffices to consider $p \in [0, \frac{1}{2}]$. By concavity,

$$\mathcal{I}_\pi(p) \geq \mathcal{I}_\pi\left(\frac{1}{2}\right) p. \quad (2.4.25)$$

Hence, in order to prove Cheeger’s isoperimetric inequality, we need only find a suitable lower bound for $\mathcal{I}_\pi(\frac{1}{2})$.

The next idea is that instead of applying the L^1 – L^∞ directly to an indicator function $\mathbb{1}_A$, we will first regularize $\mathbb{1}_A$ using the Langevin semigroup $(P_t)_{t \geq 0}$ with stationary distribution π . We start with a semigroup calculation.

Proposition 2.4.26. *Assume that the Markov semigroup $(P_t)_{t \geq 0}$ is reversible and satisfies the curvature-dimension condition $CD(0, \infty)$. For all $t > 0$ and $p \in [2, \infty]$,*

$$\|\sqrt{\Gamma(P_t f)}\|_{L^p(\pi)} \leq \frac{1}{\sqrt{2t}} \|f\|_{L^p(\pi)} \quad (2.4.27)$$

and

$$\|f - P_t f\|_{L^1(\pi)} \leq \sqrt{2t} \|\sqrt{\Gamma(f)}\|_{L^1(\pi)}. \quad (2.4.28)$$

Proof The local reverse Poincaré inequality (2.2.19) yields

$$\sqrt{\Gamma(P_t f)} \leq \frac{1}{\sqrt{2t}} \sqrt{P_t(f^2)}$$

so that

$$\left\| \sqrt{\Gamma(P_t f)} \right\|_{L^p(\pi)} \leq \frac{1}{\sqrt{2t}} \{ \mathbb{E}_\pi P_t(|f|^p) \}^{1/p} \leq \frac{1}{\sqrt{2t}} \|f\|_{L^p(\pi)}.$$

The last inequality is the dual of the $p = \infty$ case. Indeed, for g with $\|g\|_{L^\infty(\pi)} \leq 1$,

$$\partial_t \int (f - P_t f) g \, d\pi = - \int P_t \mathcal{L} f g \, d\pi = \int \Gamma(f, P_t g) \, d\pi$$

and hence, by the Cauchy–Schwarz inequality for the carré du champ ([Exercise 1.5](#)),

$$\begin{aligned} \int (f - P_t f) g \, d\pi &= \int_0^t \left(\int \Gamma(f, P_s g) \, d\pi \right) ds \leq \int_0^t \left(\int \sqrt{\Gamma(f)} \sqrt{\Gamma(P_s g)} \, d\pi \right) ds \\ &\leq \left\| \sqrt{\Gamma(f)} \right\|_{L^1(\pi)} \int_0^t \left\| \sqrt{\Gamma(P_s g)} \right\|_{L^\infty(\pi)} ds \\ &\leq \left\| \sqrt{\Gamma(f)} \right\|_{L^1(\pi)} \|g\|_{L^\infty(\pi)} \int_0^t \frac{1}{\sqrt{2s}} ds \leq \sqrt{2t} \left\| \sqrt{\Gamma(f)} \right\|_{L^1(\pi)}. \quad \square \end{aligned}$$

We are now ready to prove Milman’s theorem.

Proof of Milman’s theorem, Theorem 2.4.22 By approximating the indicator function $\mathbb{1}_A$ with a smooth function and applying the inequality (2.4.28), we can justify the bound

$$\sqrt{2t} \pi^+(A) \geq \|\mathbb{1}_A - P_t \mathbb{1}_A\|_{L^1(\pi)}.$$

Next, a calculation shows that

$$\begin{aligned} \mathbb{E}_\pi |\mathbb{1}_A - P_t \mathbb{1}_A| &= 2 \left\{ \pi(A) \pi(A^c) - \mathbb{E} \left[(\mathbb{1}_A - \pi(A)) (P_t \mathbb{1}_A - \pi(A)) \right] \right\} \\ &\geq 2 \left\{ \pi(A) \pi(A^c) - \underbrace{\|\mathbb{1}_A - \pi(A)\|_{L^\infty(\pi)} \|P_t \mathbb{1}_A - \pi(A)\|_{L^1(\pi)}}_{\leq 1} \right\}. \end{aligned}$$

From the L^1 – L^∞ Poincaré inequality and (2.4.27),

$$\|P_t \mathbb{1}_A - \pi(A)\|_{L^1(\pi)} \leq C_{1,\infty} \|\nabla P_t \mathbb{1}_A\|_{L^\infty(\pi)} \leq \frac{C_{1,\infty}}{\sqrt{2t}} \|\mathbb{1}_A\|_{L^\infty(\pi)} \leq \frac{C_{1,\infty}}{\sqrt{2t}}.$$

Hence, we have

$$\sqrt{2t} \pi^+(A) \geq 2 \left\{ \pi(A) \pi(A^c) - \frac{C_{1,\infty}}{\sqrt{2t}} \right\}.$$

Now choose $\sqrt{2t} = 2C_{1,\infty}/(\pi(A) \pi(A^c))$ to obtain

$$\pi^+(A) \geq \frac{1}{2C_{1,\infty}} \pi(A)^2 \pi(A^c)^2.$$

This inequality is not fully satisfactory, but if we take $\pi(A) = \frac{1}{2}$ then we deduce from this that $\mathcal{I}_\pi(\frac{1}{2}) \geq 1/(32C_{1,\infty})$, and from (2.4.25) we conclude. $\square$

2.4.5 Gaussian Isoperimetry

The Cheeger isoperimetric inequality asserts that for small p , $\mathcal{I}_\pi(p) \gtrsim p$. On the other hand, one can check that the Gaussian isoperimetric profile $\mathcal{I}_{\gamma_d}(p) = \phi(\Phi^{-1}(p))$ has the asymptotics $\mathcal{I}_{\gamma_d}(p) \sim p\sqrt{2\log(1/p)}$ as $p \searrow 0$.

As with the Cheeger isoperimetric inequality, isoperimetry of Gaussian type can also be captured via a functional inequality. The following result is due to [Bobkov \(1997\)](#).

Theorem 2.4.29 (Gaussian isoperimetry, functional form). *Suppose that $\pi \in \mathcal{P}(\mathbb{R}^d)$ is α -strongly log-concave for some $\alpha > 0$. Then, for all $f : \mathbb{R}^d \rightarrow [0, 1]$,*

$$\sqrt{\alpha} \mathcal{I}_{\gamma_d}(\mathbb{E}_\pi f) \leq \mathbb{E}_\pi \sqrt{\alpha \mathcal{I}_{\gamma_d}(f)^2 + \Gamma(f)}. \quad (2.4.30)$$

As this formulation suggests, [Theorem 2.4.29](#) has a proof via Markov semigroup theory, for which we refer readers to [Bakry et al. \(2014, Section 8.5.2\)](#). By converting the functional inequality back into an isoperimetric statement, one can deduce the following comparison theorem.

Theorem 2.4.31 (Gaussian isoperimetry comparison theorem). *Suppose that the measure $\pi \in \mathcal{P}(\mathbb{R}^d)$ is α -strongly log-concave for some $\alpha > 0$. Then,*

$$\mathcal{I}_\pi \geq \sqrt{\alpha} \mathcal{I}_{\gamma_d}.$$

We explore these results further in the exercises.

2.5 Riemannian Manifolds

In this section, we revisit the curvature-dimension condition. As we hinted at in [Section 2.2.1](#), the commutation relation which underlies the curvature-dimension condition captures the underlying curvature of both the ambient space and the measure. To aid the reader who is unfamiliar with Riemannian geometry, we will provide a brief review of the main concepts. In this section, we shall omit many of the proofs, as the goal is simply to acquaint the reader with the general picture in a geometric context.

2.5.1 Riemannian Geometry

Basic concepts.

We recall some of the definitions from [Section 1.3.2](#). A **Riemannian manifold** $\mathcal{M}$ is a space which is locally homeomorphic to a Euclidean space, such that at every point $p \in \mathcal{M}$ there is an associated vector space $T_p\mathcal{M}$, called the **tangent space** to $\mathcal{M}$ at p , equipped with an inner product $\langle \cdot, \cdot \rangle_p$. The tangent space $T_p\mathcal{M}$ represents the velocities of all curves passing through p . We can collect together the different tangent spaces into a single object called the **tangent bundle**,

$$T\mathcal{M} := \bigcup_{p \in \mathcal{M}} (\{p\} \times T_p\mathcal{M}).$$

The Riemannian metric $p \mapsto \langle \cdot, \cdot \rangle_p$ is required to be smooth in a suitable sense.

A smooth function $f : \mathcal{M} \rightarrow \mathbb{R}$ has a **differential** $df : T\mathcal{M} \rightarrow \mathbb{R}$, defined as follows. Given a point

$p \in \mathcal{M}$ and a tangent vector $v \in T_p\mathcal{M}$, let $(p_t)_{t \in \mathbb{R}}$ be a curve on $\mathcal{M}$ with $p_0 = p$ and with velocity v at time 0. Then, $(df)_p v := \partial_t|_{t=0} f(p_t)$. One can check that this definition does not depend on the choice of curve $(p_t)_{t \in \mathbb{R}}$ and that $(df)_p$ is a linear function on $T_p\mathcal{M}$. Note that the differential can be defined on any manifold, even if it does not have a Riemannian structure, but $(df)_p$ is not an element of $T_p\mathcal{M}$; it is an element of the dual space $T_p^*\mathcal{M}$, called the **cotangent space**. The Riemannian metric allows us to identify $(df)_p$ with an element of $T_p\mathcal{M}$: there is a unique vector $\nabla f(p) \in T_p\mathcal{M}$ such that for all $v \in T_p\mathcal{M}$, it holds that $(df)_p v = \langle \nabla f(p), v \rangle_p$. The vector $\nabla f(p)$ is called the **gradient** of f at p . The gradient depends on the choice of the metric, and we can then define the **gradient flow** of f to be a curve $(p_t)_{t \geq 0}$ such that the velocity $\dot{p}_t$ of the curve equals $-\nabla f(p_t)$ for all $t \geq 0$.

We pause to give a simple example. Suppose that $\mathcal{M} = \mathbb{R}^d$, and we pick a smooth mapping $p \mapsto A_p$ where A_p is a positive definite $d \times d$ matrix for each $p \in \mathbb{R}^d$. This induces a Riemannian metric via $\langle u, v \rangle_p := \langle u, A_p v \rangle$, where $\langle \cdot, \cdot \rangle$ (without a subscript) denotes the usual Euclidean inner product. If $(p_t)_{t \in \mathbb{R}}$ is a smooth curve in $\mathbb{R}^d$ and its usual time derivative is $(\dot{p}_t)_{t \in \mathbb{R}}$, then we know that $\partial_t f(p_t) = \langle \nabla f(p_t), \dot{p}_t \rangle = \langle A_{p_t}^{-1} \nabla f(p_t), \dot{p}_t \rangle_{p_t}$. Hence, the manifold gradient $\nabla_{\mathcal{M}} f$ is given by $\nabla_{\mathcal{M}} f(p) = A_p^{-1} \nabla f$, where ∇f is the Euclidean gradient. When $A_p = \nabla^2 \phi(p)$ is obtained as the Hessian of a mapping ϕ , then $\mathcal{M}$ is called a **Hessian manifold**.

A **vector field** on $\mathcal{M}$ is a mapping $X : \mathcal{M} \rightarrow T\mathcal{M}$ such that $X(p) \in T_p\mathcal{M}$ for all $p \in \mathcal{M}$.⁵ A single vector $v \in T_p\mathcal{M}$ can be thought of as a differential operator; for $f \in C^\infty(\mathcal{M})$, we can define the action of v on f via $v(f) := (df)_p v$. Similarly, a vector field X acts on f and produces a function $Xf : \mathcal{M} \rightarrow \mathbb{R}$, defined by $Xf(p) = X(p)f$ for all $p \in \mathcal{M}$. For example, on $\mathbb{R}^d$, a vector field can be identified with a mapping $\mathbb{R}^d \rightarrow \mathbb{R}^d$, and it differentiates functions via $Xf(p) = \langle \nabla f(p), X(p) \rangle$.

We would also like to differentiate vector fields along other vector fields, and there are two main ways of doing so. The first is called the Lie derivative, and it can be defined on any smooth manifold without the need for a Riemannian metric, and is consequently less important for our discussion. The second is the **Levi-Civita connection**, which given vector fields X and Y , outputs another vector field $\nabla_X Y$. This connection is characterized by various properties, including compatibility with the Riemannian metric: for all vector fields X, Y , and Z , we have the chain rule

$$Z\langle X, Y \rangle = \langle \nabla_Z X, Y \rangle + \langle X, \nabla_Z Y \rangle. \quad (2.5.1)$$

Here, the vector field Z is differentiating the scalar function $p \mapsto \langle X(p), Y(p) \rangle_p$. Since we do not aim to perform many Riemannian calculations here, we omit most of the other properties for simplicity. However, we mention one key fact, which is that for any smooth function f , it holds that $\nabla_{fX} Y = f \nabla_X Y$, where fX is the vector field $(fX)(p) = f(p) X(p)$. This property implies that the mapping $(X, Y) \mapsto \nabla_X Y$ is *tensorial* in its first argument, that is, $(\nabla_X Y)(p)$ only depends on the value $X(p)$ of X at p .

The tensorial property of the Levi-Civita connection allows us to compute the derivative of a vector field Y along a curve $c : \mathbb{R} \rightarrow \mathcal{M}$. Namely, for $t \in \mathbb{R}$, we can define $D_c Y(t) := (\nabla_{\dot{c}(t)} Y)(c(t))$, which makes sense because we can extend $\dot{c}$ to a vector field X on $\mathcal{M}$ and deduce that $(\nabla_X Y)(c(t))$ only depends on $X(c(t)) = \dot{c}(t)$ (and not on the choice of extension X). Then, $D_c Y$ is called the **covariant derivative** of Y along the curve c . From there, we can define the **parallel transport** of a vector $v_0 \in T_{c(0)}\mathcal{M}$ along the curve c to be the unique vector field $(v(t))_{t \in \mathbb{R}}$ defined along the curve c with $v(0) = v_0$ such that the covariant derivative vanishes: $D_c v = 0$. The parallel transport is a canonical way of identifying two different tangent spaces on $\mathcal{M}$. Due to compatibility with the metric, it has the

⁵ Geometers would write that X is a *section* of the tangent bundle.

property that if $c(0) = p$, $c(1) = q$, and $P_c v \in T_q \mathcal{M}$ denotes the parallel transport of $v \in T_p \mathcal{M}$ along c for time 1, then $P_c : T_p \mathcal{M} \rightarrow T_q \mathcal{M}$ is an *isometry*.

We already have seen the idea of a length-minimizing curve, or a **geodesic**. Recall that the Riemannian metric induces a distance on $\mathcal{M}$ via

$$d(p, q) = \inf \left\{ \int_0^1 \|\dot{\gamma}(t)\|_{\gamma(t)} dt \mid \gamma(0) = p, \gamma(1) = q \right\}.$$

If there is a minimizing constant-speed curve γ in this variational problem, we say that γ is a geodesic joining p and q . By taking the first variation of this problem, one shows that a necessary condition for γ to be a geodesic is for the covariant derivative of its velocity to vanish: $D_\gamma \dot{\gamma} = 0$. We will write this, however, with the more familiar notation $\ddot{\gamma} = 0$, which in Euclidean space means that there is zero acceleration (and hence Euclidean geodesics are straight lines). The converse is not true; if $\ddot{\gamma} = 0$, it does not imply that γ must be a shortest path between its endpoints (but it means that γ is locally a shortest path).

If $p \in \mathcal{M}$ and $v \in T_p \mathcal{M}$, then $\exp_p(v)$ is defined to be the endpoint (at time 1) of a constant-speed geodesic emanating from p with velocity v , if such a geodesic exists. In general, the exponential map may only be defined in a neighborhood of 0 on $T_p \mathcal{M}$. The logarithmic map is the inverse of the exponential map: given $q \in \mathcal{M}$, $\log_p(q)$ is the unique vector $v \in T_p \mathcal{M}$, if this is well-defined, such that $\exp_p(v) = q$.

Given a vector field X on $\mathcal{M}$, the **divergence** of X is the function $\operatorname{div} X : \mathcal{M} \rightarrow \mathbb{R}$ defined as follows: $(\operatorname{div} X)(p)$ is the trace⁶ of the linear mapping $v \mapsto (\nabla_v X)(p)$ on $T_p \mathcal{M}$. Also, for $f \in C^\infty(\mathcal{M})$, we define the **Hessian** of f at p to be the bilinear mapping $\nabla^2 f(p) : T_p \mathcal{M} \times T_p \mathcal{M} \rightarrow \mathbb{R}$ given by

$$\nabla^2 f(p)[v, w] = \langle \nabla_v \nabla f(p), w \rangle_p.$$

Even though we have not given all of the definitions precisely, we will now work through one example to give the reader the flavor of the computations. Suppose that $(p_t)_{t \in \mathbb{R}}$ is a curve on $\mathcal{M}$; then we know that $\partial_t f(p_t) = \langle \nabla f(p_t), \dot{p}_t \rangle_{p_t}$. If $g(p) := \langle \nabla f(p), \dot{p} \rangle_p$, then by definition we have $\partial_t^2 f(p_t) = \dot{p}_t(g)(p_t)$. By compatibility of the Levi-Civita connection with the metric (2.5.1), this equals $\langle \nabla_{\dot{p}_t} \nabla f(p_t), \dot{p}_t \rangle_{p_t} + \langle \nabla f(p_t), \nabla_{\dot{p}_t} \dot{p}_t \rangle_{p_t}$. The first term is $\nabla^2 f(p_t)[\dot{p}_t, \dot{p}_t]$. For the second term, if $(p_t)_{t \in \mathbb{R}}$ is a geodesic, then the term $\nabla_{\dot{p}_t} \dot{p}_t$ vanishes. This shows that it is convenient to pick geodesic curves when computing Hessians.⁷

For $f \in C^\infty(\mathcal{M})$, the **Laplacian** of f is the function $\Delta f : \mathcal{M} \rightarrow \mathbb{R}$ defined by $\Delta f := \operatorname{tr} \nabla^2 f$. In the Riemannian setting, Δ is usually called the **Laplace–Beltrami operator**. The Riemannian metric induces a **volume measure**, which we always denote via $\mathbf{m}$. Throughout, when we abuse notation to refer to the density of an absolutely continuous measure $\mu \in \mathcal{P}(\mathcal{M})$, we always refer to the density w.r.t. the volume measure, i.e., $\frac{d\mu}{d\mathbf{m}}$. We have the integration by parts formula

$$\int \Delta f g d\mathbf{m} = \int f \Delta g d\mathbf{m} = - \int \langle \nabla f, \nabla g \rangle d\mathbf{m},$$

⁶The trace is defined as usual, namely if A is a linear mapping on $T_p \mathcal{M}$, then after choosing an arbitrary orthonormal basis $e_1, \dots, e_d$ of $T_p \mathcal{M}$ (w.r.t. the Riemannian metric), we have $\operatorname{tr} A = \sum_{i=1}^d \langle e_i, A e_i \rangle_p$.

⁷When computing first-order derivatives, it is only important that the first-order behavior of the curve is correct (i.e., the curve has the correct tangent vector). When computing second-order derivatives, it should come at no surprise that the second-order behavior of the curve begins to matter.

Incidentally, if $\nabla f(p_t) = 0$, i.e., we are at a *stationary point*, then the second term vanishes regardless of the curve p . Hence, the Hessian of f can be defined on any smooth manifold without the need for a Riemannian metric, provided that we restrict ourselves to stationary points of f . This observation is used heavily in Morse theory.

provided that there are no boundary terms.

Curvature.

For a two-dimensional surface, it is easier to define the notion of curvature: one has the **Gaussian curvature**, which associates to each point $p \in \mathcal{M}$ a single number $K(p) \in \mathbb{R}$. It is the product of the two principal curvatures at p . The celebrated *Theorema Egregium* (“remarkable theorem”) of Gauss asserts that the Gaussian curvature is unchanged under local isometries, i.e., the Gaussian curvature is intrinsic to the surface. (In contrast, there are other *extrinsic* notions of curvature, such as the mean curvature, which rely on the embedding of the manifold in Euclidean space.)

In higher dimensions, we are not so fortunate and it requires much more geometric information to fully capture the idea of curvature. In fact, at each point $p \in \mathcal{M}$, we associate to it a 4-tensor, called the **Riemann curvature tensor**. It is defined as follows: given vector fields W, X, Y , and Z ,

$$\text{Riem}(W, X, Y, Z) := \langle \nabla_X \nabla_W Y - \nabla_W \nabla_X Y + \nabla_{[W, X]} Y, Z \rangle.$$

Here, $[W, X]$ is the **Lie bracket** of W and X , which is the vector field U defined as the commutator: $Uf := WXf - XWf$. This tensor is obviously an unwieldy object. Nevertheless, we may begin to get a handle on it by observing that at its core, it measures the lack of commutativity of certain differential operators, which we stated was the basis for curvature in Section 2.2.1. On Euclidean space, it vanishes: $\text{Riem} = 0$. Also, the Riemann curvature tensor is fully determined by the **sectional curvatures** of $\mathcal{M}$: given a two-dimensional subspace S of $T_p \mathcal{M}$, the sectional curvature of S can be defined as the Gaussian curvature of the two-dimensional surface obtained by following geodesics with directions in S . Thus, we can view the Riemann curvature tensor as collecting together all of the curvature information from two-dimensional slices.

Luckily, the Riemann curvature tensor contains information that is too detailed for our purposes. With an eye toward probabilistic applications, we focus mainly on properties such as the distortion of volumes of balls along geodesics, which only requires looking at certain averages of the Riemann curvature. More specifically, for $u, v \in T_p \mathcal{M}$, let

$$\text{Ric}_p(u, v) := \text{tr} \text{Riem}(u, \cdot, v, \cdot).$$

The tensor Ric is called the **Ricci curvature tensor**. It is a powerful fact that many useful geometric and probabilistic consequences, such as diameter bounds and functional inequalities, are consequences of lower bounds on the Ricci curvature.

We also mention that one can further take the trace of the Ricci curvature tensor to arrive at a single scalar function, known as the **scalar curvature**, but we shall not use it in this book.

Diffusions on manifolds.

Recall that on $\mathbb{R}^d$, the generator of the standard Brownian motion is $\frac{1}{2} \Delta$, where Δ is the Laplacian operator. On a manifold $\mathcal{M}$, we define standard Brownian motion $(B_t)_{t \geq 0}$ to be the unique $\mathcal{M}$ -valued stochastic process with generator $\frac{1}{2} \Delta$, where Δ is now the Laplace–Beltrami operator. This means that for all smooth functions $f : \mathcal{M} \rightarrow \mathbb{R}$, we require $t \mapsto f(B_t) - f(B_0) - \int_0^t \frac{1}{2} \Delta f(B_s) ds$ to be a local martingale; see [Hsu \(2002\)](#).

More generally, a stochastic process $(Z_t)_{t \geq 0}$ has generator $\mathcal{L}$ if for all smooth functions $f : \mathcal{M} \rightarrow \mathbb{R}$, the process $t \mapsto f(Z_t) - f(Z_0) - \int_0^t \mathcal{L} f(Z_s) ds$ is a local martingale. When the generator is $\mathcal{L} f = \Delta f - \langle \nabla V, \nabla f \rangle$ for a smooth function $V : \mathcal{M} \rightarrow \mathbb{R}$, this corresponds to a Langevin diffusion on the manifold. We informally write $dZ_t = -\nabla V(Z_t) dt + \sqrt{2} dB_t$, although the “+” symbol has to be

interpreted carefully. Under some assumptions, the stationary distribution π of the Langevin diffusion has density $\pi \propto \exp(-V)$ w.r.t. the volume measure $\mathfrak{m}$.

Under appropriate assumptions on ∇V , the existence and uniqueness of the diffusion process on the manifold can be proven, e.g., via embedding the manifold in Euclidean space and using similar arguments as in Section 1.1.3.

Optimal transport on Riemannian manifolds.

We conclude this section by discussing how the optimal transport problem can be generalized to Riemannian manifolds. Recall from [Exercise 1.10](#) that the optimal transport problem can be posed with other costs; in particular, we take the cost to be $c(x, y) = \mathbf{d}(x, y)^2$, where $\mathbf{d}$ is the distance induced by the Riemannian metric. Suppose, for simplicity, that $\mathcal{M}$ is compact and that μ is absolutely continuous (w.r.t. the volume measure). Then, there is a unique optimal transport map T from μ to ν , which is of the form $T(x) = \exp_x(\nabla\psi(x))$, where $-\psi$ is $\mathbf{d}^2/2$ -concave.

Moreover, there is a formal Riemannian structure on $\mathcal{P}_{2,\text{ac}}(\mathcal{M})$. We can formally define the tangent space at μ to be

$$T_\mu \mathcal{P}_{2,\text{ac}}(\mathcal{M}) := \overline{\{\nabla\psi \mid \psi \in C^\infty(\mathcal{M})\}}^{L^2(\mu)},$$

equipped with the norm $\|\nabla\psi\|_\mu := \sqrt{\int \|\nabla\psi\|^2 d\mu}$. Also, curves of measures are again characterized by the continuity equation

$$\partial_t \mu_t + \operatorname{div}(\mu_t v_t) = 0,$$

where the equation is to be interpreted in a weak sense: for any test function $\varphi \in C^\infty(\mathcal{M})$, for a.e. t , it holds that

$$\partial_t \int \varphi d\mu_t = \int \langle \nabla\varphi, v_t \rangle d\mu_t.$$

In short, aside from new technicalities introduced in the Riemannian setting (such as the presence of a cut locus⁸), most of the facts familiar to us from the Euclidean setting continue to hold when generalized appropriately. We refer to [Villani \(2009b\)](#) for more details.

2.5.2 Geometry of Markov Semigroups

We now indicate how Markov semigroup proofs can be extended to the setting of a weighted Riemannian manifold $\mathcal{M}$ with a reference measure π which admits a density $\pi \propto \exp(-V)$ w.r.t. the volume measure $\mathfrak{m}$.

Consider the Langevin diffusion on $\mathcal{M}$ with generator $\mathcal{L}$ given by

$$\mathcal{L}f := \Delta f - \langle \nabla V, \nabla f \rangle.$$

As before, we can compute the carré du champ to be

$$\Gamma(f) = \|\nabla f\|^2.$$

⁸Loosely speaking, the presence of a cut locus means that there are multiple minimizing geodesics connecting two points. Think for instance of the two poles of a sphere.

For the iterated carré du champ,

$$\begin{aligned}\Gamma_2(f) &= \frac{1}{2} \{ \mathcal{L}(\|\nabla f\|^2) - 2 \langle \nabla f, \nabla \mathcal{L} f \rangle \} \\ &= \frac{1}{2} \{ \Delta(\|\nabla f\|^2) - 2 \langle \nabla f, \nabla \Delta f \rangle \} + \langle \nabla f, \nabla^2 V \nabla f \rangle.\end{aligned}$$

Unlike in Section 2.2.1, however, we now have to apply the Bochner identity

$$\frac{1}{2} \Delta(\|\nabla f\|^2) = \langle \nabla f, \nabla \Delta f \rangle + \|\nabla^2 f\|_{\text{HS}}^2 + \text{Ric}(\nabla f, \nabla f) \quad (2.5.2)$$

which shows that

$$\Gamma_2(f) = \|\nabla^2 f\|_{\text{HS}}^2 + \langle \nabla f, (\text{Ric} + \nabla^2 V) \nabla f \rangle.$$

Observe that in this formula, the curvature of the ambient space and the curvature of the measure are placed on an equal footing through the tensor $\text{Ric} + \nabla^2 V$. If $\text{Ric} + \nabla^2 V \geq \alpha$, in the sense that $\text{Ric}(X, X) + \langle X, \nabla^2 V X \rangle \geq \alpha \|X\|^2$ for any vector field X on $\mathcal{M}$, then the curvature-dimension condition $\Gamma_2 \geq \alpha \Gamma$ holds. Since the proof of the Bakry–Émery theorem (Theorem 1.2.30) only relied on the $\text{CD}(\alpha, \infty)$ condition (together with calculus rules for the Markov semigroup, such as the chain rule), the theorem continues to hold in the setting of weighted Riemannian manifolds.

Actually, we can refine the condition further as follows. If $\dim \mathcal{M} = d$, then

$$\|\nabla^2 f\|_{\text{HS}}^2 \geq \frac{1}{d} (\text{tr } \nabla^2 f)^2 = \frac{1}{d} (\Delta f)^2.$$

This observation motivates the following definition.

Definition 2.5.3. A Markov semigroup is said to satisfy the **curvature-dimension condition** with curvature lower bound α and dimension bound d , denoted $\text{CD}(\alpha, d)$, if for all functions f ,

$$\Gamma_2(f) \geq \alpha \Gamma(f) + \frac{1}{d} (\mathcal{L} f)^2. \quad (2.5.4)$$

As the name suggests, the following theorem holds.

Theorem 2.5.5. Let $\mathcal{M}$ be a complete Riemannian manifold with volume measure $\mathfrak{m}$, and let $\alpha > 0$, $d \geq 1$. Consider the Markov semigroup associated with Brownian motion on $\mathcal{M}$, with generator Δ . Then, the following two statements are equivalent.

- 1 $\text{CD}(\alpha, d)$ holds.
- 2 $\text{Ric} \geq \alpha$ and $\dim \mathcal{M} \leq d$.

As an example, one can show that the unit sphere $\mathbb{S}^d$ satisfies $\text{Ric} = d - 1$, so that the $\text{CD}(d - 1, d)$ condition holds. Then, using the Bakry–Émery theorem (Theorem 1.2.30), or by using Markov semigroup calculus to prove that the curvature-dimension condition implies Bobkov’s functional form of the Gaussian isoperimetric inequality (Theorem 2.4.29; see Bakry et al., 2014, Corollary 8.5.4), one can now deduce results such as concentration on the sphere (Theorem 2.4.6).

Besides providing an abstract framework for deriving functional inequalities, it is worth noting that the condition (2.5.4) no longer makes any mention of the ambient space except through the Markov semigroup $(P_t)_{t \geq 0}$ and its associated operators $\mathcal{L}$, Γ , and Γ_2 . This has led to a line of research

investigating to what extent we can study the intrinsic geometry associated with a Markov semigroup. Although we do not survey the literature here, we show one illustrative example to give the flavor of the results. First, one shows that the $\text{CD}(\alpha, d)$ condition implies a Sobolev inequality.

Theorem 2.5.6 (Sobolev inequality). *Consider a diffusion Markov semigroup satisfying the $\text{CD}(\alpha, d)$ condition for some $\alpha > 0$ and $d > 2$. Then, for all $p \in [1, \frac{2d}{d-2}]$ and all functions f ,*

$$\frac{1}{p-2} \left\{ \left(\int |f|^p d\pi \right)^{2/p} - \int f^2 d\pi \right\} \leq \frac{d-1}{\alpha d} \int \Gamma(f) d\pi. \quad (2.5.7)$$

From this Sobolev inequality, one can then deduce a diameter bound for the Markov semigroup. Here, the diameter is defined as follows:

$$\text{diam}((P_t)_{t \geq 0}) := \sup_{x, y \in \mathcal{X}} \{ \pi\text{-ess sup } |f(x) - f(y)| \mid \|\Gamma(f)\|_{L^\infty(\pi)} \leq 1 \}.$$

Theorem 2.5.8 (Diameter bound). *Suppose that the Markov semigroup $(P_t)_{t \geq 0}$ satisfies the Sobolev inequality (2.5.7). Then,*

$$\text{diam}((P_t)_{t \geq 0}) \leq \pi \sqrt{\frac{d-1}{\alpha}}.$$

The diameter bound is sharp, as it is attained by the sphere, and together with [Theorem 2.5.6](#) it recovers the classical **Bonnet–Myers diameter bound** from Riemannian geometry. Other geometric results obtained in this fashion include volume growth comparison results and heat kernel bounds.

2.5.3 Displacement Convexity of Entropy

The other perspective with which we can encode geometry is the optimal transport perspective. Namely, in [Section 1.4](#), we informally argued that in the Euclidean context, the α -strong convexity of the KL divergence $\text{KL}(\cdot \parallel \pi)$ on $(\mathcal{P}_2(\mathbb{R}^d), W_2)$ is equivalent to the α -strong convexity of the potential V . At this stage, it is a perhaps expected, although still remarkable, fact that on a general weighted Riemannian manifold $\mathcal{M}$, the α -strong convexity of $\text{KL}(\cdot \parallel \pi)$ on $(\mathcal{P}_2(\mathcal{M}), W_2)$ is equivalent to the $\text{CD}(\alpha, \infty)$ condition, which in turn is equivalent to $\text{Ric} + \nabla^2 V \geq \alpha$.

Theorem 2.5.9. *Let $\mathcal{M}$ be a complete connected Riemannian manifold and let $\pi \propto \exp(-V)$ be a smooth probability density w.r.t. the volume measure. Then, the following are equivalent:*

- 1 *It holds that $\text{Ric} + \nabla^2 V \geq \alpha$.*
- 2 *For all measures $\mu_0, \mu_1 \in \mathcal{P}_2(\mathcal{M})$, there exists a constant-speed geodesic $(\mu_t)_{t \in [0, 1]}$ joining μ_0 to μ_1 such that for all $t \in [0, 1]$,*

$$\text{KL}(\mu_t \parallel \pi) \leq (1-t) \text{KL}(\mu_0 \parallel \pi) + t \text{KL}(\mu_1 \parallel \pi) - \frac{\alpha t(1-t)}{2} W_2^2(\mu_0, \mu_1). \quad (2.5.10)$$

A key observation is that the condition (2.5.10) still makes sense when we replace $\mathcal{M}$ with a space that admits a notion of geodesics, i.e., there is no need for a smooth structure *a priori*. Similarly,

there are ways to formulate the general $\text{CD}(\alpha, d)$ condition via displacement convexity, but they are considerably more complicated, and we omit them for simplicity. This observation is the starting place for the influential **Lott–Sturm–Villani theory of synthetic Ricci curvature lower bounds** (Sturm, 2006a,b; Lott and Villani, 2009). See the supplement for further context and details.

2.6 Discrete Space and Time

Up until this point, we have been focusing on continuous-time Markov processes on a continuous state space. In this section, we give a few pointers on what may break down in discrete space or discrete time. Our treatment here is far from comprehensive.

Discrete space.

For Markov processes on a discrete space, we can still define the Markov semigroup, generator, carré du champ, and Dirichlet form. The main difference is that the carré du champ is now a finite difference operator, rather than a differential operator, and consequently it fails to satisfy a chain rule.

Crucially, this difference manifests itself for the log-Sobolev inequality, which we have written in this chapter as

$$\text{ent}_\pi(f^2) \leq 2C \mathcal{E}(f) \quad \text{for all } f. \quad (2.6.1)$$

On the other hand, recall from [Theorem 1.2.26](#) (which still holds for discrete state spaces) that the exponential decay of the KL divergence is equivalent to the inequality

$$\text{ent}_\pi(f) \leq \frac{C}{2} \mathcal{E}(f, \log f) \quad \text{for all } f \geq 0. \quad (2.6.2)$$

When the carré du champ satisfies a chain rule, then (2.6.1) and (2.6.2) are equivalent, but in general the first inequality (2.6.1) is strictly stronger.

Lemma 2.6.3. *The inequality (2.6.1) implies inequality (2.6.2).*

See [Exercise 2.23](#). The first inequality (2.6.1) is often simply called the *log Sobolev inequality*, whereas the second inequality (2.6.2) is called a **modified log-Sobolev inequality (MLSI)**. In many cases, the log-Sobolev inequality is too strong in that it does not hold with a good constant C ; hence, the modified log-Sobolev inequality is often the more appropriate inequality for the discrete setting.

Discrete time.

Similarly, for discrete-time Markov chains we can no longer use semigroup calculus, although the basic principles of Poincaré inequalities (spectral gap inequalities) and modified log-Sobolev inequalities can be adapted to this setting. In addition, there are new techniques based on the notion of *conductance*, which extend the concept of isoperimetry as discussed in [Section 2.4](#). As we shall need to study discrete-time Markov chains in detail for sampling algorithms, we defer a fuller discussion of this theory to [Chapter 7](#).

Discrete curvature.

Inspired by the geometric connections in [Section 2.5](#), many researchers have attempted to define notions of curvature on discrete spaces. We do not attempt to survey this literature here, but we give a few pointers to the literature.

Ollivier (Ollivier, 2007, 2009) introduced the following notion of curvature.

Definition 2.6.4. A metric space $(\mathcal{X}, d)$ equipped with a Markov kernel P is said to have **coarse Ricci curvature** bounded below by $\kappa \in [0, 1]$ if for all $x, y \in \mathcal{X}$,

$$W_1(P(x, \cdot), P(y, \cdot)) \leq (1 - \kappa) d(x, y). \quad (2.6.5)$$

In other words, the Markov chain with kernel P is a W_1 contraction. The definition is motivated by the following observation: on a d -dimensional Riemannian manifold with $\text{Ric} \geq \alpha$, let $P(x, \cdot)$ be the uniform measure on $B(x, \varepsilon)$. Then, provided that $d(x, y) = O(\varepsilon)$, it holds that

$$W_1(P(x, \cdot), P(y, \cdot)) \leq \left(1 - \frac{\alpha \varepsilon^2}{2(d+2)} + O(\varepsilon^3)\right) d(x, y).$$

A lower bound on the coarse Ricci curvature is often too strong of an assumption for the purpose of studying mixing times of Markov chains, although there are refinements in Ollivier (2007, 2009). However, when a lower bound on the coarse Ricci curvature holds, then it implies a number of useful consequences, such as concentration estimates and functional inequalities. We mention the following result in particular (Exercise 2.25).

Theorem 2.6.6 (Ollivier (2009)). *Suppose that P is a reversible Markov kernel on a metric space $(\mathcal{X}, d)$, and that P has coarse Ricci curvature bounded below by κ . Then, P satisfies a Poincaré inequality with constant at most $1/\kappa$.*

Refer to Chapter 7 for a precise definition of the Poincaré inequality used here. One interesting feature of Theorem 2.6.6 is that the coarse Ricci curvature depends on the choice of metric d , but the Poincaré inequality conclusion does not, which leaves flexibility in the choice of d .

Other approaches for studying the curvature of discrete Markov processes include: studying the displacement convexity of entropy (using different interpolating curves rather than W_2 geodesics) (Ollivier and Villani, 2012; Gozlan et al., 2014; Léonard, 2017); using ideas from Bakry–Émery theory (Klartag et al., 2016; Fathi and Shu, 2018); and defining a modified W_2 distance for which the Markov process becomes a gradient flow of the KL divergence (Maas, 2011; Erbar and Maas, 2012; Mielke, 2013).

We emphasize that although we only described the coarse Ricci curvature approach in any detail, there is not a single approach which supersedes the others in the discrete setting. Each approach has its own merits and shortcomings.

Bibliographical Notes

The monographs Bakry et al. (2014); van Handel (2016) are excellent sources to learn more about Markov semigroup theory.

The Monge–Ampère equation introduced in Exercise 2.1, being a fully non-linear PDE, is notoriously difficult to study. See Villani (2003, Chapter 4) for an overview of rigorous results on the Monge–Ampère equation, including the celebrated regularity theory of Caffarelli. The proofs of Proposition 2.1.1 and Exercise 2.3 are taken from Cordero-Erausquin (2017). One might wonder whether a “log-Sobolev” version of the Brascamp–Lieb inequality holds, but the answer is unfortunately negative (Bobkov and Ledoux, 2000).

The general idea used in the proof of Lemma 2.2.7 is attributed to Hörmander (2007). In the proof

of Lemma 2.2.7, we assumed the solvability of the Poisson equation; this can be avoided via a density argument, see Cordero-Erausquin et al. (2004); Barthe and Cordero-Erausquin (2013). The proof of the dimensional Brascamp–Lieb inequality in Exercise 2.4 is taken from Bolley et al. (2018), and the bound on $\text{var}_\pi V$ therein was used in Chewi (2021) to show that the entropic barrier is an optimal self-concordant barrier. Also, Exercise 2.5 is from Helffer and Sjöstrand (1994), and weighted variants can be found in Arnaudon et al. (2018); whereas Exercise 2.6 is from Klartag (2023). Finally, we caution the reader that the Brascamp–Lieb inequality in Theorem 2.2.9 should not be confused with another family of inequalities, which are unfortunately also known as Brascamp–Lieb inequalities and described in, e.g., Villani (2003, Section 6.3).

The convergence in Rényi divergence of the Langevin diffusion was obtained earlier, under stronger assumptions in Cao et al. (2019). A natural question to ask is whether there are functional inequalities that interpolate between the Poincaré and log-Sobolev inequalities, which imply intermediate rates of convergence for the Langevin diffusion. One answer is given by the family of **Latala–Oleszkiewicz inequalities (LOI)** (Latala and Oleszkiewicz, 2000). The convergence of the Langevin diffusion under an LO inequality is given in Chewi et al. (2025a). One can also consider super Poincaré inequalities (Wang, 2000), or variants of Sobolev inequalities (Chafaï, 2004).

For further works investigating functional inequalities in the low temperature regime, see Kinoshita and Suzuki (2022); Li and Erdogdu (2023); Chen and Sridharan (2025); Gong et al. (2025).

For the Lipschitz perturbation result (Proposition 2.3.4), there was a trick attributed to L. Miclo (see Bardet et al., 2018, Lemma 2.1) which rewrites the Lipschitz perturbation as a bounded one and applies Holley–Stroock; unfortunately, this did not yield a dimension-free result. The first dimension-independent bound was due to Brigati and Pedrotti (2024) based on constructing a Lipschitz transport map from the Gaussian; we give this approach as Exercise 3.9. Our proof, which yields a sharper constant, is inspired by Davies et al. (2026, Lemma 6.6).

The Prékopa–Leindler inequality given in Exercise 2.11 can be used to deduce other functional inequalities, such as the log-Sobolev and Bregman transport inequalities; see Bobkov and Ledoux (2000); Gentil (2008). It was subsequently generalized to Riemannian manifolds in Cordero-Erausquin et al. (2001). For more on preservation of log-concavity, see Saumard and Wellner (2014).

Exercise 2.14 essentially contains the main results of Chen et al. (2021a) (actually the paper assumes a slightly weaker condition than (2.E.3), namely that the p -th moment of the chi-squared divergence is bounded for some $p > 1$, but this is handled with the same arguments as in Exercise 2.14). Earlier works on functional inequalities for mixtures include Chafaï and Malrieu (2010); Zimmermann (2013, 2016); Bardet et al. (2018); Schlichting (2019).

There are many treatments on concentration of measure, e.g., Ledoux (2001); Boucheron et al. (2013); Bakry et al. (2014); van Handel (2016); Vershynin (2018). The Poincaré case of Theorem 2.4.3 is from Bobkov (1999). The monographs Bobkov and Houdré (1997); Ledoux (2001); Bakry et al. (2014) are excellent sources to learn about isoperimetry. The exposition of the functional form of Cheeger’s inequality (Theorem 2.4.18) as well as Milman’s theorem (Theorem 2.4.22) were inspired by the treatment in Alonso-Gutiérrez and Bastero (2015). It would be hard to survey the various developments on this subject here, but we would like to mention a few nice additions to the story. First, as we saw in Theorem 2.4.18 and Proposition 2.4.21, isoperimetric inequalities are typically stronger than their functional inequality counterparts, and often strictly so. In order to obtain inequalities involving sets which are *equivalent* to, say, the Poincaré and log-Sobolev inequalities, one should turn toward *measure capacity inequalities*, for which we refer the reader to Bakry et al. (2014, Chapter 8).

Also, more refined “two-level” isoperimetric inequalities have been pioneered in [Talagrand \(1991\)](#), which has applications in its own right.

From [Exercise 2.15](#), the Harnack inequality was first established in [Wang \(1997\)](#), and the log-Harnack inequality proof is from [Bobkov et al. \(2001\)](#). These inequalities are variants of parabolic Harnack inequalities from the PDE literature, and they are notable for holding with dimension-free constants. Moreover, whereas parabolic Harnack inequalities typically compare the semigroup at two different times, the inequalities from [Exercise 2.15](#) avoids this by considering the commutation of the Markov semigroup with strictly convex functions such as the power function $(\cdot)^p$. There is a huge literature on these inequalities which we cannot survey here, but we refer to F.-Y. Wang’s monographs for details ([Wang, 2013, 2014](#)). The convexity principle is taken from [Altschuler and Chewi \(2024b\)](#), which also contains further discussion. Further proofs of the inequalities in [Exercise 2.17](#) are given as [Exercise 3.4](#) and [Exercise 4.7](#).

The Gaussian isoperimetric inequality is due to [Sudakov and Cirel’son \(1974\)](#) and [Borell \(1975\)](#). It has since been extended and refined in various ways, e.g., in the context of noise stability ([Borell, 1985](#); [Isaksson and Mossel, 2012](#); [Eldan, 2015](#); [Mossel and Neeman, 2015](#); [Kindler et al., 2018](#)).

Section 2.5 draws upon many resources, which we list here: [do Carmo \(1992\)](#) for Riemannian geometry; [Hsu \(2002\)](#) for diffusions on manifolds; [Villani \(2009b\)](#) for optimal transport on manifolds, including synthetic Ricci curvature lower bounds; and [Ledoux \(2000\)](#); [Bakry et al. \(2014\)](#) for the geometry of Markov semigroups. The curvature-dimension condition and its equivalences have been explored in a vast number of works, e.g., [Sturm \(2006a,b\)](#); [Lott and Villani \(2009\)](#); [Wang \(2011\)](#).

Although the MLSI is generally weaker than the LSI, the works [Salez et al. \(2023\)](#); [Salez and Youssef \(2025\)](#) have obtained conditions under which the LSI constant can be bounded by the MLSI constant, at least for finite-state reversible chains.

Exercises

Overview of the Inequalities

⊇ [Exercise 2.1](#) (Linearization of the Monge–Ampère equation)

In general, when $\mu, \nu \in \mathcal{P}_{2,\text{ac}}(\mathbb{R}^d)$ have smooth densities and $\nabla\varphi$ denotes the optimal transport map from μ to ν , then from the change of variables formula we expect

$$\frac{\mu}{\nu \circ \nabla\varphi} = \det \nabla^2\varphi.$$

This is known as the **Monge–Ampère equation**. It is a non-linear PDE in the variable φ , which is a convex function (by Brenier’s theorem, [Theorem 1.3.8](#)). In this exercise, we linearize the Monge–Ampère equation to gain insight into the infinitesimal behavior of optimal transport.

Let μ be a probability measure on $\mathbb{R}^d$ with a smooth density, and let $f \in C_c^\infty(\mathbb{R}^d)$ satisfy $\int f d\mu = 0$. Let $\mu_\varepsilon := (1 + \varepsilon f)\mu$, and let $\nabla\varphi_\varepsilon$ denote the optimal transport map from μ to μ_ε . Assuming that $\varphi_\varepsilon(x) = \frac{\|x\|^2}{2} + \varepsilon u(x) + o(\varepsilon)$ for some function $u : \mathbb{R}^d \rightarrow \mathbb{R}$, perform an expansion of the Monge–Ampère equation in ε and argue that u satisfies the following linear PDE, known as the **Poisson equation**:

$$-\mathcal{L}u = f, \quad \text{where} \quad \mathcal{L}u := \Delta u - \left\langle \nabla \log \frac{1}{\mu}, \nabla u \right\rangle.$$

Note that $\mathcal{L}$ is the generator of the Langevin diffusion with stationary distribution μ (see [Example 1.2.4](#)).

Use this to formally argue that

$$\lim_{\varepsilon \searrow 0} \frac{1}{\varepsilon^2} W_2^2(\mu, (1 + \varepsilon f) \mu) = \int \|\nabla u\|^2 d\mu = \int u (-\mathcal{L}) u d\mu = \int f (-\mathcal{L})^{-1} f d\mu.$$

Here, $\int \|\nabla u\|^2 d\mu = \int u (-\mathcal{L}) u d\mu$ is the squared **Sobolev norm** $\|u\|_{\dot{H}^1(\mu)}^2$, where the dot is used to distinguish this from the usual Sobolev norm $\|u\|_{H^1(\mu)}^2 = \|u\|_{L^2(\mu)}^2 + \|u\|_{\dot{H}^1(\mu)}^2$. Similarly, $\int f (-\mathcal{L})^{-1} f d\mu$ is the squared **inverse Sobolev norm** $\|f\|_{\dot{H}^{-1}(\mu)}^2$. Therefore, the linearization result shows that $W_2^2(\mu, \nu) \sim \|\mu - \nu\|_{\dot{H}^{-1}(\mu)}^2$ as $\nu \rightarrow \mu$.

Using the linearization (1.E.1) of the KL divergence from [Exercise 1.7](#), deduce that the $T_2(C)$ inequality implies

$$C \int f^2 d\mu \geq \int f (-\mathcal{L})^{-1} f d\mu.$$

In light of the spectral gap interpretation of the Poincaré inequality, why does the above inequality suggest that $T_2(C)$ implies $PI(C)$?

The astute reader should also work out how the Poisson equation can be obtained starting with the continuity equation (1.3.18).

Proofs via Markov Semigroup Theory

⊇ [Exercise 2.2](#) (Curvature-dimension condition)

Verify the commutation identity (2.2.4) and the formula (2.2.5) for the iterated carré du champ.

⊇ [Exercise 2.3](#) (Bregman transport inequality)

Let $\nabla\varphi$ denote the optimal transport map from π to μ , so that the Monge–Ampère equation holds (see [Exercise 2.1](#)):

$$\frac{\pi}{\mu \circ \nabla\varphi} = \det \nabla^2 \varphi.$$

Take logarithms of both sides of this equation and integrate w.r.t. π to prove the Bregman transport inequality ([Theorem 2.2.12](#)). Then, by applying [Proposition 2.1.1](#), give another proof of the Brascamp–Lieb inequality ([Theorem 2.2.9](#)).

⊇ [Exercise 2.4](#) (Dimensional improvement of the Brascamp–Lieb inequality)

In finite-dimensional space, one can improve upon the Brascamp–Lieb inequality ([Theorem 2.2.9](#)) by subtracting a non-negative term from the right-hand side. There are different ways to do this, but in this exercise we explore an approach which utilizes the extra term $\|\nabla^2 u\|_{\text{HS}}^2$ in the iterated carré du champ operator.

- 1 Let $\pi \propto \exp(-V)$, where as before we assume that V is twice continuously differentiable and strictly convex. Let f satisfy $\mathbb{E}_\pi f = 0$, and consider another function u (not necessarily the solution to $-\mathcal{L}u = f$). Show that

$$\mathbb{E}_\pi [f^2] \leq \mathbb{E}_\pi [(f + \mathcal{L}u)^2] + \mathbb{E}_\pi \langle \nabla f, (\nabla^2 V)^{-1} \nabla f \rangle - \mathbb{E}_\pi [\|\nabla^2 u\|_{\text{HS}}^2].$$

- 2 Prove that $\mathbb{E}_\pi [\|\nabla^2 u\|_{\text{HS}}^2] \geq d^{-1} (\mathbb{E}_\pi \Delta u)^2$, and that

$$\mathbb{E}_\pi \Delta u = \text{cov}_\pi(f, V) - \mathbb{E}_\pi [(f + \mathcal{L}u) V].$$

- 3 Choose u to solve $-\mathcal{L}u = f + \lambda (V - \mathbb{E}_\pi V)$ for some $\lambda \geq 0$ and substitute this into the previous parts. Optimize over λ and prove that

$$\mathrm{var}_\pi f \leq \mathbb{E}_\pi \langle \nabla f, (\nabla^2 V)^{-1} \nabla f \rangle - \frac{\mathrm{cov}_\pi(f, V)^2}{d - \mathrm{var}_\pi V}.$$

- 4 In particular, deduce that

$$\mathrm{var}_\pi V \leq \frac{d \mathbb{E}_\pi \langle \nabla V, (\nabla^2 V)^{-1} \nabla V \rangle}{d + \mathbb{E}_\pi \langle \nabla V, (\nabla^2 V)^{-1} \nabla V \rangle} \leq d.$$

≥ **Exercise 2.5** (Helffer–Sjöstrand identity)

Let $\mathcal{L}$ denote the generator corresponding to $\pi \propto \exp(-V)$, $\nabla^2 V > 0$. Heuristically derive the following *identity*: for all $f, g : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\mathrm{cov}_\pi(f, g) = \mathbb{E}_\pi \langle \nabla f, (-\mathcal{L} + \nabla^2 V)^{-1} \nabla g \rangle.$$

Note that since $-\mathcal{L} \geq 0$, this implies the Brascamp–Lieb inequality (Theorem 2.2.9)!

Hint: Let $(\cdot)^*$ denote the adjoint in $L^2(\pi)$. It suffices to prove that $\nabla^* (-\mathcal{L} + \nabla^2 V)^{-1} \nabla$ is the orthogonal projection onto $1^\perp$. Suppose that $L^2(\pi)$ admits a basis of eigenfunctions for $-\mathcal{L}$, and let $u : \mathbb{R}^d \rightarrow \mathbb{R}$ be such an eigenfunction with $-\mathcal{L}u = \lambda u$, $\lambda > 0$. Study the effect of the operator on u , recalling (1.2.18) and (2.2.4).

≥ **Exercise 2.6** (Improved Lichnerowicz inequality)

Suppose that π is α -strongly log-concave. The goal of this exercise is to prove the following improved Lichnerowicz inequality:

$$C_{\mathrm{Pl}}(\pi) \leq \sqrt{\frac{\|\mathrm{cov}_\pi\|_{\mathrm{op}}}{\alpha}}.$$

Note that this interpolates between the Bakry–Émery result $C_{\mathrm{Pl}}(\pi) \leq 1/\alpha$ and the KLS conjecture $C_{\mathrm{Pl}}(\pi) \lesssim \|\mathrm{cov}_\pi\|_{\mathrm{op}}$.

- 1 Let u be an eigenfunction corresponding to the spectral gap of $-\mathcal{L}$. Let $\lambda := 1/C_{\mathrm{Pl}}(\pi)$. By writing $\lambda^2 = \int (\mathcal{L}u)^2 d\pi$, bringing in the iterated carré du champ (as in the proof of Lemma 2.2.7), and applying strong log-concavity and the Poincaré inequality, show that $\|\int \nabla u d\pi\|^2 \geq \alpha$.
- 2 Next, write $\|\int \nabla u d\pi\| = \sup_{\theta \in \mathbb{R}^d, \|\theta\|=1} \int \langle \nabla u, \theta \rangle d\pi$. Use the integration by parts formula and Cauchy–Schwarz to show that $\|\int \nabla u d\pi\|^2 \leq \lambda^2 \|\mathrm{cov}_\pi\|_{\mathrm{op}}$. Conclude the proof.

≥ **Exercise 2.7** (Local Poincaré inequality)

In Theorem 2.2.17, prove that $\mathrm{CD}(\alpha, \infty)$ implies (2.2.19), and that a Taylor expansion of either (2.2.18) or (2.2.19) implies back $\mathrm{CD}(\alpha, \infty)$.

Hint: To establish (2.2.19), use (2.2.20).

≥ **Exercise 2.8** (Diffusion chain rule)

Verify the equivalence between (2.2.14) and (2.2.15).

≥ **Exercise 2.9** (Monotonicity of Rényi divergences in the order)

Prove that for $1 \leq q \leq q' \leq \infty$, $\mathcal{R}_q \leq \mathcal{R}_{q'}$.

▷ **Exercise 2.10** (Hypercontractivity)

Let $(P_t)_{t \geq 0}$ be a Markov semigroup with stationary distribution π . Assume that the semigroup either satisfies the diffusion chain rule, or that it is reversible. Show that the LSI in the form

$$\text{ent}_\pi(f^2) \leq 2C_{\text{LSI}} \mathcal{E}(f)$$

for all f is equivalent to the following **hypercontractivity** statement: for all functions f , $t \geq 0$, and $p \geq 1$, if we set $p(t) := 1 + (p - 1) \exp(2t/C_{\text{LSI}})$, then

$$\|P_t f\|_{L^{p(t)}(\pi)} \leq \|f\|_{L^p(\pi)}.$$

This is a strengthening of the fact that the semigroup is a contraction on any $L^p(\pi)$ space and shows that in fact the semigroup maps $L^p(\pi)$ into $L^{p'}(\pi)$ for some $p' > p$.

Hint: Differentiate $t \mapsto \log \|P_t f\|_{L^{p(t)}(\pi)}$. In the case where the semigroup is assumed to be reversible (but does not satisfy the diffusion chain rule), one needs to show the following inequality: for any non-negative function h and any $q > 1$, it holds that $\mathcal{E}(h^{q-1}, h) \geq \frac{4(q-1)}{q^2} \mathcal{E}(h^{q/2})$. Show that this follows from the numerical inequality $(b^{q-1} - a^{q-1})(b - a) \geq \frac{4(q-1)}{q^2} (b^{q/2} - a^{q/2})^2$ for all $0 \leq a \leq b$.

▷ **Exercise 2.11** (Prékopa–Leindler inequality)

In this exercise, we introduce another important functional inequality, known as the **Prékopa–Leindler inequality**.

- 1 Let π be α -strongly log-concave, let $t \in [0, 1]$, and let $f, g, h : \mathbb{R}^d \rightarrow \mathbb{R}_{>0}$ be three functions such that for all $x, y \in \mathbb{R}^d$,

$$\log h((1-t)x + ty) \geq (1-t) \log f(x) + t \log g(y) - \frac{\alpha}{2} t(1-t) \|x - y\|^2.$$

Prove that

$$\log \int h \, d\pi \geq (1-t) \log \int f \, d\pi + t \log \int g \, d\pi. \quad (2.E.1)$$

Hint: Let μ_0, μ_1 achieve equality in the Donsker–Varadhan variational principle (Theorem 1.5.7), so that

$$\begin{aligned} \log \mathbb{E}_\pi f &= \mathbb{E}_{\mu_0} \log f - \text{KL}(\mu_0 \parallel \pi), \\ \log \mathbb{E}_\pi g &= \mathbb{E}_{\mu_1} \log g - \text{KL}(\mu_1 \parallel \pi). \end{aligned}$$

Let μ_t be along the Wasserstein geodesic from μ_0 to μ_1 . Apply the Donsker–Varadhan principle again, together with the assumption on f, g, h as well as strong convexity of the KL divergence, in order to lower bound $\log \mathbb{E}_\pi h$.

- 2 The inequality (2.E.1) continues to hold when π is replaced by Lebesgue measure, if we set $\alpha = 0$ in the assumption.⁹ Use this to prove that if π is a log-concave measure over $\mathbb{R}^{d_1+d_2}$, then the marginal π^1 of π on $\mathbb{R}^{d_1}$ is also log-concave.

Hint: Partition elements of $\mathbb{R}^{d_1+d_2}$ as (x, y) . Apply the Prékopa–Leindler inequality on $\mathbb{R}^{d_2}$ with $f := \pi(x_0, \cdot)$, $g := \pi(x_1, \cdot)$, and $h := \pi((1-t)x_0 + tx_1, \cdot)$.

⁹For example, one could first consider the case when f, g, h are compactly supported, take π to be the uniform distribution over a ball $B(0, R)$, and take $R \rightarrow \infty$.

- 3 We now aim to generalize the fact in the previous part. Suppose that π is a density on $\mathbb{R}^{d_1+d_2}$ such that $(x, y) \mapsto \log \pi(x, y) + \frac{1}{2} \langle (x, y), \Sigma^{-1} (x, y) \rangle$ is concave, where

$$\Sigma = \begin{bmatrix} \Sigma_{1,1} & \Sigma_{1,2} \\ \Sigma_{2,1} & \Sigma_{2,2} \end{bmatrix} > 0.$$

Prove that for the marginal π^1 of π on $\mathbb{R}^{d_1}$, $x \mapsto \log \pi^1(x) + \frac{1}{2} \langle x, (\Sigma_{1,1})^{-1} x \rangle$ is concave. (Use the result in the previous part.)

- 4 Show that if π is α -strongly log-concave, then the convolution $\pi * \text{normal}(0, tI_d)$ is $\alpha/(1 + \alpha t)$ -strongly log-concave.
- 5 Show that the Prékopa–Leindler inequality for the Lebesgue measure is equivalent to the **Brunn–Minkowski inequality**: for compact sets $A, B \subseteq \mathbb{R}^d$,

$$\text{vol}((1-t)A + tB) \geq \text{vol}(A)^{1-t} \text{vol}(B)^t. \quad (2.E.2)$$

By scaling A and B and choosing t , show that (2.E.2) can be upgraded to

$$\text{vol}(A + B)^{1/d} \geq \text{vol}(A)^{1/d} + \text{vol}(B)^{1/d}.$$

Operations Preserving Functional Inequalities

⊢ Exercise 2.12 (Variational principle for entropies)

Let $\phi : \mathbb{R}_{>0} \rightarrow \mathbb{R}$ be a convex differentiable function and let D_ϕ denote the associated Bregman divergence (cf. Definition 2.2.11). For any positive random variable X with $\mathbb{E}|\phi(X)| < \infty$ and any $t > 0$, prove that $\mathbb{E} D_\phi(X, t) - \mathbb{E} D_\phi(X, \mathbb{E} X) = D_\phi(\mathbb{E} X, t)$, and deduce that $\mathbb{E} \phi(X) - \phi(\mathbb{E} X) = \inf_{t>0} \mathbb{E} D_\phi(X, t)$. Use this to prove the variational principle for the entropy used in the proof of Holley–Stroock perturbation (Proposition 2.3.1).

⊢ Exercise 2.13 (Transport inequality in one dimension)

Let π be the standard Gaussian on $\mathbb{R}$, and let $\mu \ll \pi$. In one dimension, the optimal transport map T from π to μ is the monotone rearrangement that satisfies, for each $x \in \mathbb{R}$, $\mu((-\infty, T(x)]) = \pi((-\infty, x])$.

- 1 Differentiate this relation to obtain a formula for $\frac{d\mu}{d\pi}(T(x))$.
- 2 Substitute this into the KL divergence $\text{KL}(\mu \parallel \pi) = \int \log \frac{d\mu}{d\pi}(T(x)) \pi(dx)$ and use the inequality $t - 1 - \log t \geq 0$ for all $t > 0$ in order to prove the Gaussian T_2 inequality in one dimension. Deduce the Gaussian T_2 inequality in general dimension via a tensorization argument.
- 3 Can you generalize this calculation to a density $\pi \propto \exp(-V)$ on $\mathbb{R}^d$, where V is smooth and α -strongly convex for some $\alpha > 0$?

⊢ Exercise 2.14 (Generalizing the LSI for mixtures)

In this exercise, we generalize the log-Sobolev inequality for mixtures (Proposition 2.3.14).

- 1 First, show that Example 2.3.15 is sharp up to universal constants as follows. Consider the case when $\mu = \frac{1}{2}(\delta_{-R} + \delta_{+R})$ on $\mathbb{R}$, so that μP is a mixture of two Gaussians. Construct a test function $f : \mathbb{R} \rightarrow \mathbb{R}$ for the Poincaré inequality which shows that $C_{\text{PI}}(\mu P) \gtrsim R^2 \exp(\Omega(R^2/\sigma^2))$ if $R/\sigma \gtrsim 1$.
- 2 Next, consider the setting of Proposition 2.3.9 except that we replace the assumption (2.3.10) with the weaker condition

$$C_{\chi^2,2} := \sqrt{\mathbb{E}[\chi^2(P_X \parallel P_{X'})^2]} < \infty, \quad (2.E.3)$$

where $X, X' \stackrel{\text{i.i.d.}}{\sim} \mu$. Now, rather than writing $\text{var } \mathbb{E}_{P_X} f = \frac{1}{2} \mathbb{E}[|\mathbb{E}_{P_X} f - \mathbb{E}_{P_{X'}} f|^2]$, instead write $\text{var } \mathbb{E}_{P_X} f = \mathbb{E}[|\mathbb{E}_{P_X} f - \mathbb{E}_{\mu P} f|^2]$. By bounding this quantity in two different ways deduce that

$$\begin{aligned} \text{var } \mathbb{E}_{P_X} f &\leq \mathbb{E} \min\{(\text{var}_{P_X} f) \chi^2(\mu P \| P_X), (\text{var}_{\mu P} f) \chi^2(P_X \| \mu P)\} \\ &\leq \mathbb{E} \sqrt{(\text{var}_{P_X} f) \chi^2(\mu P \| P_X) (\text{var}_{\mu P} f) \chi^2(P_X \| \mu P)}. \end{aligned}$$

Use this to prove that a Poincaré inequality holds for μP , and give an upper bound on $C_{\text{PI}}(\mu P)$.

- 3 Now consider the setting of [Proposition 2.3.14](#) except that we again assume the weaker condition (2.E.3). Previously, we bounded

$$\mathbb{E}[\mathbb{E}_{P_X}(f^2) \log(1 + \chi^2(P_X \| P_{X'}))] \leq \mathbb{E}_{\mu P}(f^2) \log(1 + C_{\chi^2}),$$

which relies on L^1 – L^∞ duality. This time, we want to use duality between “ $L \log L$ ” and “ $\exp L$ ”. Namely, use the variational principle for the entropy ([Lemma 2.3.5](#)) to prove that for a suitable constant $C > 0$ (depending on $C_{\chi^2, 2}$),

$$2 \mathbb{E}[\mathbb{E}_{P_X}(f^2) \{\log(1 + \chi^2(P_X \| P_{X'})) - C\}] \leq \text{ent } \mathbb{E}_{P_X}(f^2).$$

Use this to prove that a log-Sobolev inequality holds for μP , and give an upper bound on $C_{\text{LSI}}(\mu P)$.

- 4 Consider [Example 2.3.15](#) again, except instead of assuming that μ is supported on $\mathbf{B}(0, R)$, we assume that μ has sub-Gaussian tails:

$$\iint \exp \frac{\|x - x'\|^2}{\sigma_{\text{SG}}^2} \mu(dx) \mu(dx') \leq C_{\text{SG}}.$$

Prove that if $\sigma \gtrsim \sigma_{\text{SG}}$ for a sufficiently large implied constant, then the Gaussian mixture μP satisfies a log-Sobolev inequality, and give an upper bound on $C_{\text{LSI}}(\mu P)$. Also, show how this can recover the result of [Example 2.3.15](#).

The next group of exercises introduces Harnack and reverse transport inequalities.

▷ **Exercise 2.15 (Harnack and reverse transport inequalities)**

Let P be a Markov kernel over a space $\mathcal{X}$, and let $x, y \in \mathcal{X}$.

- 1 Use Donsker–Varadhan duality to show that

$$Pf(x) \leq C_{\log}(x, y) + \log P(\exp f)(y) \quad \text{for all bounded } f : \mathcal{X} \rightarrow \mathbb{R}, \quad (2.E.4)$$

if and only if

$$\text{KL}(\delta_x P \| \delta_y P) \leq C_{\log}(x, y). \quad (2.E.5)$$

- 2 Similarly, let $p, q \in (1, \infty)$ be a pair of Hölder conjugate exponents, and use duality to show that

$$Pf(x) \leq C_p(x, y) P(f^p)(y)^{1/p} \quad \text{for all bounded } f : \mathcal{X} \rightarrow \mathbb{R}_+, \quad (2.E.6)$$

if and only if

$$\mathcal{R}_q(\delta_x P \| \delta_y P) \leq \frac{q}{q-1} \log C_p(x, y). \quad (2.E.7)$$

We refer to (2.E.4) as a **log-Harnack inequality**, (2.E.6) as a **Harnack inequality**, and (2.E.5) and (2.E.7) as **reverse transport inequalities**.

- 3 Show that if (2.E.6) holds for every $p > 1$ with $C_p = \exp(\frac{C}{p-1})$, then (2.E.4) holds with $C_{\log} \leq C$. *Hint:* Let $p \rightarrow \infty$.

▷ **Exercise 2.16** (Convexity principle)

Let P be a Markov kernel over $\mathcal{X}$.

- 1 Suppose that $\text{KL}(\delta_x P \parallel \delta_y P) \leq \rho(x, y)$ for all $x, y \in \mathcal{X}$. Using the joint convexity of the KL divergence (Theorem 1.5.6), prove that

$$\text{KL}(\mu P \parallel \nu P) \leq \inf_{\gamma \in \mathcal{C}(\mu, \nu)} \int \rho(x, y) \gamma(dx, dy).$$

- 2 Suppose that $\mathcal{R}_q(\delta_x P \parallel \delta_y P) \leq \rho(x, y)$ for all $x, y \in \mathcal{X}$. Although $\mathcal{R}_q$ is not jointly convex for $q > 1$, it is a monotonic transformation of an f -divergence. Use this to give an upper bound on $\mathcal{R}_q(\mu P \parallel \nu P)$ in terms of an optimal transport problem between μ and ν .

▷ **Exercise 2.17** (Harnack inequalities via semigroup arguments)

Let $(P_t)_{t \geq 0}$ be a Markov semigroup over $\mathbb{R}^d$, satisfying the $\text{CD}(\alpha, \infty)$ condition with $\alpha \in \mathbb{R}$.

- 1 Let $x, y \in \mathbb{R}^d$ and let $(x_s)_{s \in [0, t]}$ be a curve joining x to y . Differentiate $P_s(\log P_{t-s} f)(x_s)$ w.r.t. s and apply the diffusion chain rule and the gradient bound (Theorem 2.2.6). By optimizing over the choice of curves $(x_s)_{s \in [0, t]}$, prove the log-Harnack inequality

$$P_t(\log f)(x) \leq \log P_t f(y) + \frac{\alpha \|x - y\|^2}{2(\exp(2\alpha t) - 1)}.$$

By Exercise 2.15 and Exercise 2.16, it implies that for all $\mu, \nu \in \mathcal{P}_2(\mathbb{R}^d)$,

$$\text{KL}(\mu P_t \parallel \nu P_t) \leq \frac{\alpha W_2^2(\mu, \nu)}{2(\exp(2\alpha t) - 1)}. \quad (2.E.8)$$

Note that this strengthens (1.4.13) in that it allows the second argument to be νP_t for any $\nu \in \mathcal{P}_2(\mathbb{R}^d)$.

- 2 Find an analogous Markov semigroup argument to prove that

$$(P_t f(x))^p \leq P_t(f^p)(y) \exp\left(\frac{\alpha p \|x - y\|^2}{2(p-1)(\exp(2\alpha t) - 1)}\right), \quad (2.E.9)$$

and hence

$$\mathcal{R}_q(\delta_x P_t \parallel \delta_y P_t) \leq \frac{\alpha q \|x - y\|^2}{2(\exp(2\alpha t) - 1)}. \quad (2.E.10)$$

- 3 Here, we consider an alternative argument. Starting from the local log-Sobolev inequality in the form (2.2.22), establish (2.E.9) by differentiating $s \mapsto \frac{p}{\xi_s} \log P_t(f^{\xi_s})(x_s)$, where $\xi_s := 1 + (p-1)s$. Argue that this reasoning can be reversed, so that (2.E.9) implies back the local log-Sobolev inequality, and hence $\text{CD}(\alpha, \infty)$.

Concentration of Measure and Isoperimetry

▷ **Exercise 2.18** (Isoperimetry on the sphere)

Prove Theorem 2.4.6 from the spherical isoperimetric inequality (Theorem 2.4.5). To do so, use the fact that the measure of $\mathbf{B}(x_0, r)$ is

$$\sigma_d(\mathbf{B}(x_0, r)) = \frac{\int_0^r (\sin \theta)^{d-1} d\theta}{\int_0^\pi (\sin \theta)^{d-1} d\theta},$$

or prove this fact yourself. It is also acceptable to establish a weaker bound of the form $\alpha_{\sigma_d}(\varepsilon) \leq C \exp(-c d \varepsilon^2)$ for universal constants $c, C > 0$.

▷ **Exercise 2.19** (Gaussian isoperimetry)

Consider the Gaussian isoperimetric inequality in Theorem 2.4.7.

- 1 In the spirit of Theorem 2.4.18, show that the functional inequality (2.4.30) is equivalent to an isoperimetric statement. Consequently, deduce the comparison theorem (Theorem 2.4.31) from Theorem 2.4.29.
- 2 Show that the functional form of the Gaussian isoperimetric inequality in (2.4.30) is preserved (up to constants) under Lipschitz mappings (in other words, prove the analogue of Proposition 2.3.3 for (2.4.30)).

▷ **Exercise 2.20** (Gaussian isoperimetry implies LSI)

In a similar spirit to Exercise 1.7, show that the Gaussian isoperimetric inequality (2.4.30) implies the log-Sobolev inequality.

▷ **Exercise 2.21** (Isoperimetry and optimal transport in dimension one)

Let $\mu, \pi \in \mathcal{P}_{2,ac}(\mathbb{R})$, and let T denote the optimal transport map from π to μ .

- 1 Show that $T = F_\mu^{-1} \circ F_\pi$, where F_μ, F_π denote the CDFs of μ and π respectively.
- 2 Use the Monge–Ampère equation (Exercise 2.1) to deduce that T is L -Lipschitz if and only if $\mu \circ F_\mu^{-1} \geq L^{-1} \pi \circ F_\pi^{-1}$, i.e., the isoperimetric profile of μ is lower bounded by the isoperimetric profile of π .
- 3 When $\pi = \text{normal}(0, I_d)$, the existence of a Lipschitz transport map from π to μ implies that μ satisfies a Gaussian isoperimetric inequality, by Exercise 2.19. Use the preceding part to deduce the converse of this statement holds in dimension one. What happens when π is the Laplace distribution with density $\pi(x) = \frac{1}{2} \exp(-|x|)$?

Riemannian Manifolds

▷ **Exercise 2.22** (Lichnerowicz inequality)

Under the $\text{CD}(\alpha, d)$ condition (2.5.4), the spectral gap estimate for $-\mathcal{L}$ can be sharpened to $\lambda_1(-\mathcal{L}) \geq \alpha d / (d - 1)$, an estimate that is attributed to Lichnerowicz. Prove this as follows: assume that f is such that $-\mathcal{L}f = \lambda f$. Show that $\lambda \int \Gamma(f) d\pi = \int \Gamma_2(f) d\pi$. Apply $\text{CD}(\alpha, d)$ and deduce that $\lambda \geq \alpha d / (d - 1)$.

Discrete Space and Time

▷ **Exercise 2.23** (LSI implies MLSI)

Prove Lemma 2.6.3.

Hint: Prove that $4(\sqrt{a} - \sqrt{b})^2 = (\int_a^b t^{-1/2} dt)^2 \leq (\log a - \log b)(a - b)$ for all $a, b > 0$.

▷ **Exercise 2.24** (Coarse Ricci curvature on graphs)

Show that if (2.6.5) holds for all $x, y \in \mathcal{X}$, then $W_1(\mu P, \nu P) \leq (1 - \kappa) W_1(\mu, \nu)$ for all $\mu, \nu \in \mathcal{P}_1(\mathcal{X})$. Furthermore, show that when d is the shortest path metric on a graph, it suffices to check (2.6.5) for neighbors, that is, for pairs (x, y) with $d(x, y) = 1$.

▷ **Exercise 2.25** (Coarse Ricci curvature implies PI)

Suppose that $\mathcal{X}$ has bounded diameter and that P satisfies (2.6.5). Let $f : \mathcal{X} \rightarrow \mathbb{R}$ be Lipschitz and satisfy $\int f d\pi = 0$, where π is the stationary distribution for P . Prove that $\limsup_{n \rightarrow \infty} (\text{var}_\pi P^n f)^{1/n} \leq (1 - \kappa)^2$, and deduce that the operator norm of P restricted to the subspace $\mathcal{V} \subseteq L^2(\pi)$ consisting of

Lipschitz functions with zero mean is at most $1 - \kappa$. Next, using the fact that $\mathcal{V}$ is dense in $L^2(\pi)$, establish [Theorem 2.6.6](#). (See [Section 7.4](#) for background on discrete-time chains.)

In this chapter, we further expand our toolbox of stochastic analysis. Namely, we introduce Girsanov's theorem, which furnishes a formula for the Radon–Nikodym derivative of the laws of two SDEs w.r.t. each other, and we discuss the time reversal of an SDE. In order to highlight the flexibility and power of these ideas, we then study some interesting applications, not all of which are directly relevant to log-concave sampling but nevertheless fit within the broader themes of this book.

3.1 Quadratic Variation

We now take a more general view of the ideas that led to the construction of the Itô integral as well as Itô's formula ([Theorem 1.1.19](#)).

Finite variation vs. quadratic variation.

As a first step toward understanding the difficulties we faced when constructing the Itô integral, we recall the classical condition under which it is possible to integrate a continuous process $(\eta_t)_{t \in [0, T]}$ against another continuous process $(A_t)_{t \in [0, T]}$ to form the integral $\int_0^T \eta_t dA_t$. This condition states that the process A is of **finite variation**. This means that for any partition $0 = t_0 < t_1 < \dots < t_n = T$ of $[0, T]$, if we define the **mesh** of the partition to be

$$\text{mesh}(t_i : i = 0, 1, \dots, n) := \max_{i \in [n]} |t_i - t_{i-1}|,$$

then it holds that

$$\lim_{\text{mesh}(t_i : i=0,1,\dots,n) \searrow 0} \sum_{i=1}^n |A_{t_i} - A_{t_{i-1}}| < \infty.$$

The above limit is called the **total variation** of A on $[0, T]$. Under this condition, there is a signed measure μ_A such that for all $t \in [0, T]$, we have $\mu_A([0, t]) = A_t - A_0$. Moreover, we can define a norm $\|\cdot\|_{\text{TV}}$ on the space of signed measures, called the **total variation norm**, for which $\|\mu_A\|_{\text{TV}}$

equals the total variation of A as defined above.¹ In this case, we can simply define the integral $\int_{[0,T]} \eta_t dA_t := \int_{[0,T]} \eta_t \mu_A(dt)$. Note that if $t \mapsto A_t$ is continuously differentiable, then the total variation of A equals $\int_{[0,T]} |\dot{A}_t| dt$, and the integral becomes $\int_{[0,T]} \eta_t dA_t = \int_{[0,T]} \eta_t \dot{A}_t dt$.

Hence, the condition that A is of finite variation is enough to develop a satisfactory theory of integration. The difficulty, however, is that Brownian motion is *not* of finite variation. To see this, take $t_i := iT/n$ for $i = 0, 1, \dots, n$, so that the mesh of the partition is T/n . Since $B_{t_i} - B_{t_{i-1}} \sim \text{normal}(0, T/n)$, we expect (heuristically) that

$$\liminf_{n \rightarrow \infty} \sum_{i=1}^n \underbrace{|B_{t_i} - B_{t_{i-1}}|}_{\asymp \sqrt{T/n}} \gtrsim \liminf_{n \rightarrow \infty} n \cdot \sqrt{\frac{T}{n}} = \infty.$$

On the other hand, if we change the definition slightly, then we expect (heuristically) that

$$\limsup_{n \rightarrow \infty} \sum_{i=1}^n \underbrace{|B_{t_i} - B_{t_{i-1}}|^2}_{\asymp T/n} \lesssim \limsup_{n \rightarrow \infty} n \cdot \frac{T}{n} < \infty.$$

We say that Brownian motion has finite **quadratic variation**. We will show in fact that the above limit is well-defined and non-trivial in the sense of convergence in probability.

More generally, for a process of the form

$$X_t = X_0 + \int_0^t b_s ds + \int_0^t \sigma_s dB_s, \quad t \in [0, T],$$

the second term is a process of finite variation (provided that $\int_{[0,T]} |b_t| dt < \infty$ almost surely), whereas the third term requires consideration of quadratic variation.

Definition of the quadratic variation.

More formally, we have the following theorem.

Theorem 3.1.1 (Quadratic variation). *Let $(M_t)_{t \in [0,T]}$ be a continuous local martingale; then, there is an a.s. unique, adapted, continuous, increasing process $t \mapsto [M, M]_t$ such that $[M, M]_0 = 0$ and $t \mapsto M_t^2 - [M, M]_t$ is a continuous local martingale. Also, suppose that for each $n \in \mathbb{N}^+$, $\{t_i\}_{i=0}^n$ is a partition of $[0, t]$, with mesh tending to zero as $n \rightarrow \infty$. Then,*

$$[M, M]_t = \lim_{n \rightarrow \infty} \sum_{i=1}^n (M_{t_i} - M_{t_{i-1}})^2 \quad \text{in probability.}$$

Definition 3.1.2. The process $[M, M]$ of Theorem 3.1.1 is called the **quadratic variation** of M .

Note that Theorem 3.1.1 provides another interpretation of the quadratic variation: when we compose the martingale M with the convex function $(\cdot)^2$, then M^2 is no longer a martingale.² The quadratic variation is the process that we can subtract from M^2 in order to make it a martingale.

¹Indeed, the notation $\|\mu - \nu\|_{TV}$ for the total variation distance between μ and ν is in accordance with this more general notion of a norm on the space of signed measures, up to a factor of 2 in the conventions.

²It is, however, a submartingale.

We will not prove [Theorem 3.1.1](#) in full generality. However, we will verify that the quadratic variation of one-dimensional Brownian motion $(B_t)_{t \in [0, T]}$ is $[B, B]_T = T$, which gives an idea of the general result. By independence of the Brownian increments,

$$\begin{aligned} \mathbb{E} \left[\left| \sum_{i=1}^n \{ (B_{t_i} - B_{t_{i-1}})^2 - (t_i - t_{i-1}) \} \right|^2 \right] &= \sum_{i=1}^n \mathbb{E} [| (B_{t_i} - B_{t_{i-1}})^2 - (t_i - t_{i-1}) |^2] \\ &\leq \sum_{i=1}^n \mathbb{E} [(B_{t_i} - B_{t_{i-1}})^4] = 3 \sum_{i=1}^n (t_i - t_{i-1})^2 \\ &\leq 3 \text{mesh}(t_i : i = 0, 1, \dots, n) \underbrace{\sum_{i=1}^n (t_i - t_{i-1})}_{=T} \rightarrow 0. \end{aligned}$$

Hence, $\sum_{i=1}^n (B_{t_i} - B_{t_{i-1}})^2 \xrightarrow{\mathbb{P}} T$ as $n \rightarrow \infty$. We also know that $t \mapsto B_t^2 - t$ is a martingale (see, e.g., [Exercise 1.5](#)).

Semimartingales.

We often consider solutions to SDEs with non-zero drift coefficients, which therefore are not continuous local martingales. To accommodate this addition, we consider the following definition.

Definition 3.1.3. A process $(X_t)_{t \in [0, T]}$ is called a **continuous semimartingale** if we can write $X = A + M$, where A is an adapted process of finite variation with $A_0 = 0$ and M is a continuous local martingale.

The decomposition $X = A + M$ is then unique. Indeed, suppose that $X = \tilde{A} + \tilde{M}$ for another finite variation process $\tilde{A}$ (with $\tilde{A}_0 = 0$) and a continuous local martingale $\tilde{M}$. Then, from $\Delta := M - \tilde{M} = \tilde{A} - A$, we deduce that Δ is both a continuous local martingale and a process of finite variation. Since Δ is of finite variation,

$$\sum_{i=1}^n (\Delta_{t_i} - \Delta_{t_{i-1}})^2 \leq \left(\max_{i \in [n]} |\Delta_{t_i} - \Delta_{t_{i-1}}| \right) \underbrace{\sum_{i=1}^n |\Delta_{t_i} - \Delta_{t_{i-1}}|}_{\text{bounded as } n \rightarrow \infty} \rightarrow 0$$

as the mesh size tends to zero. This shows that $[\Delta, \Delta] = 0$, and thus Δ^2 is a continuous local martingale. If we knew that Δ^2 were a genuine martingale, then together with $\Delta_0 = 0$ it would imply that $\Delta = 0$, establishing uniqueness of the semimartingale decomposition. We omit the localization argument required to finish the proof.

We can also define the quadratic variation of the semimartingale X as $[X, X] := [M, M]$. To see

why this makes sense, observe that

$$\begin{aligned} & \left| \sum_{i=1}^n (X_{t_i} - X_{t_{i-1}})^2 - \sum_{i=1}^n (M_{t_i} - M_{t_{i-1}})^2 \right| \\ &= \left| \sum_{i=1}^n (A_{t_i} - A_{t_{i-1}})^2 + 2 \sum_{i=1}^n (A_{t_i} - A_{t_{i-1}})(M_{t_i} - M_{t_{i-1}}) \right| \\ &\leq \sum_{i=1}^n (A_{t_i} - A_{t_{i-1}})^2 + 2 \sqrt{\left(\sum_{i=1}^n (A_{t_i} - A_{t_{i-1}})^2 \right) \left(\sum_{i=1}^n (M_{t_i} - M_{t_{i-1}})^2 \right)}, \end{aligned}$$

which tends to zero using the same argument as in the uniqueness of the semimartingale decomposition: finite variation processes have zero quadratic variation.

The bracket of two semimartingales.

Given two semimartingales X and Y , we define their bracket via polarization:

$$[X, Y] := \frac{1}{2} ([X + Y, X + Y] - [X, X] - [Y, Y]).$$

Equivalently, if $X = A_X + M_X$ and $Y = A_Y + M_Y$ are the respective decompositions, then $[X, Y] = [M_X, M_Y]$. The following theorem gives a concrete way of computing the bracket for processes driven by Brownian motion.

Theorem 3.1.4 (Bracket of processes driven by Brownian motion). *Suppose that X and Y are $\mathbb{R}^d$ -valued processes with*

$$\begin{aligned} dX_t &= b_t^X dt + \sigma_t^X dB_t, \\ dY_t &= b_t^Y dt + \sigma_t^Y dB_t, \end{aligned}$$

where we assume $\int_{[0,T]} \|b_t^X\| dt$, $\int_{[0,T]} \|b_t^Y\| dt$, $\int_{[0,T]} \|\sigma_t^X\|_{\text{HS}}^2 dt$, and $\int_{[0,T]} \|\sigma_t^Y\|_{\text{HS}}^2 dt$ are all finite almost surely. Then, X and Y are continuous semimartingales, and

$$[X, Y]_t = \int_0^t \sigma_s^X \otimes \sigma_s^Y ds, \quad \text{for } t \in [0, T].$$

Itô's formula revisited.

Finally, we conclude this section by revisiting Itô's formula (Theorem 1.1.19) using our new calculus.

Theorem 3.1.5 (Itô's formula revisited). *Let X be an $\mathbb{R}^d$ -valued semimartingale, and write $X = (X^1, \dots, X^d)$. Let $f \in C^2(\mathbb{R}^d)$. Then, $f(X)$ is also a semimartingale, and*

$$f(X_t) = f(X_0) + \sum_{i=1}^d \int_0^t \partial_i f(X_s) dX_s^i + \frac{1}{2} \sum_{i,j=1}^d \int_0^t \partial_i \partial_j f(X_s) d[X^i, X^j]_s.$$

If we interpret $[X, X]$ as the matrix whose (i, j) -entry is $[X^i, X^j]$, then this can be written in matrix

notation as

$$f(X_t) = f(X_0) + \int_0^t \langle \nabla f(X_s), dX_s \rangle + \frac{1}{2} \int_0^t \langle \nabla^2 f(X_s), d[X, X]_s \rangle. \quad (3.1.6)$$

For d -dimensional standard Brownian motion $(B_t)_{t \in [0, T]}$, we have $[B, B]_t = tI_d$, so that $d[B, B]_t = I_d dt$ and we recover the original statement of Itô's formula in [Theorem 1.1.19](#). The point is that the quadratic variation is a convenient way of streamlining Itô calculations, as it formalizes the idea that only the Brownian motion part of a process contributes in the second-order term in Itô's formula.

3.2 Change of Measure in Path Space

In this section, we begin to investigate measures on path space, for which it is convenient to adopt the following canonical setup. Let $(B_t)_{t \in [0, T]}$ be standard Brownian motion, and let $\mathbf{W}$ denote its law, the **Wiener measure**. Recall that in probability theory, all of our random variables are defined on an underlying probability space $(\Omega, \mathcal{F}, \mathbb{P})$. Since we can take this probability space to be whatever we wish, as long as it is sufficiently rich, we may as well take $\Omega = C([0, T]; \mathbb{R}^d)$ to be the space of continuous paths with $\mathcal{F}$ being the Borel σ -algebra, equipped with the Wiener measure $\mathbb{P} = \mathbf{W}$. Then, for $\omega \in \Omega$, the Brownian motion at time t simply becomes the evaluation functional, $B_t(\omega) = \omega_t$.

3.2.1 The Cameron–Martin Theorem

Our goal now is to understand when two measures $\mathbf{P}, \mathbf{Q}$ on the path space $C([0, T])$ are absolutely continuous w.r.t. each other, and if so, to write down a formula for the Radon–Nikodym derivative $\frac{d\mathbf{P}}{d\mathbf{Q}}$. The final result, known as *Girsanov's theorem*, will be used for the analysis of sampling algorithms later in this book, and more broadly it is an indispensable tool for stochastic analysis.

Before reaching this goal, however, it may be helpful to provide some mathematical context. We begin with the following question. For a curve $h \in C([0, T])$, the translation operator $T_h : C([0, T]) \rightarrow C([0, T])$ is defined simply by the mapping $\omega \mapsto \omega + h$. What happens to the Wiener measure under translations?

More generally, we can define the translation operator T_h on any Banach space $\mathcal{B}$, with $h \in \mathcal{B}$. If $\mathcal{B} = \mathbb{R}^d$ and μ is the Lebesgue measure, then we know that μ is invariant under translations, in the sense that $(T_h)_\# \mu = \mu$ for all $h \in \mathbb{R}^d$, and that the Lebesgue measure is the unique measure with this property up to rescaling. However, as soon as we move to infinite dimensions, a classical result of analysis states that there is *no* non-trivial locally finite measure μ which is invariant under translations, i.e., infinite-dimensional Lebesgue measure does not exist. This makes it difficult to decide upon a “canonical” reference measure for infinite-dimensional analysis.

Although invariance is impossible, we can at least ask for *quasi-invariance*: does there exist μ such that $(T_h)_\# \mu \ll \mu$ for all $h \in \mathcal{B}$? For example, the standard Gaussian measure is quasi-invariant on $\mathbb{R}^d$. In infinite dimensions, the answer is still *no*; in particular, there is no infinite-dimensional standard Gaussian. To understand this point more concretely, suppose that $\mathcal{B} = \mathcal{H}$ is actually a Hilbert space, and let $(e_k)_{k \in \mathbb{N}}$ be an orthonormal basis. An obvious attempt to build “the standard Gaussian measure on $\mathcal{H}$ ” is to take an i.i.d. sequence $(\xi_k)_{k \in \mathbb{N}}$ of standard Gaussians on $\mathbb{R}$, and to take μ to be the law of $\sum_{k \in \mathbb{N}} \xi_k e_k$. However, since $\|\sum_{k=0}^N \xi_k e_k\|^2 = \sum_{k=0}^N \xi_k^2$, standard probability theory tells us that almost surely, this is not a convergent sum in $\mathcal{H}$.

Despite this obstruction, we will now see that the Wiener measure behaves in some sense like a

standard Gaussian measure on a Hilbert space! The theorem below begins by precisely characterizing the set of h for which $(T_h)_\# \mathbf{W} \ll \mathbf{W}$.

Theorem 3.2.1 (Cameron–Martin). *Let $\mathbf{W}$ be the Wiener measure on $C([0, T])$. Then, $(T_h)_\# \mathbf{W} \ll \mathbf{W}$ if and only if*

$$h \in \mathcal{H} := \left\{ h \in C([0, T]) \mid h(0) = 0, \int_0^T \|\dot{h}(t)\|^2 dt < \infty \right\}. \quad (3.2.2)$$

If this holds, then

$$\frac{d(T_h)_\# \mathbf{W}}{d\mathbf{W}}(\omega) = \exp\left(\int_0^T \langle \dot{h}(t), d\omega_t \rangle - \frac{1}{2} \int_0^T \|\dot{h}(t)\|^2 dt\right). \quad (3.2.3)$$

*Here, $\mathcal{H}$ is called the **Cameron–Martin space** associated with Brownian motion.*

The Cameron–Martin theorem will be subsumed by Girsanov’s theorem, so we will not prove it here. Instead, we will focus on its interpretation.

Interpretation of the Cameron–Martin theorem.

To interpret (3.2.3), let γ denote the standard Gaussian measure on $\mathbb{R}^d$ and let $h \in \mathbb{R}^d$. Then, on $\mathbb{R}^d$, we have the formula

$$\frac{d(T_h)_\# \gamma}{d\gamma}(x) = \exp\left(\langle h, x \rangle - \frac{1}{2} \|h\|^2\right).$$

This bears a striking resemblance to (3.2.3). Namely, for $h_0, h_1 \in \mathcal{H}$, let us define the inner product $\langle h_0, h_1 \rangle_{\mathcal{H}} := \int_0^T \langle \dot{h}_0(t), \dot{h}_1(t) \rangle dt$. If we interpret the stochastic integral $\int_0^T \langle \dot{h}(t), d\omega_t \rangle$ as $\langle h, \omega \rangle_{\mathcal{H}}$, then the density ratio in (3.2.3) behaves as if

$$d\mathbf{W}(\omega) \text{ “}\propto\text{” } \exp\left(-\frac{1}{2} \|\omega\|_{\mathcal{H}}^2\right) d\omega. \quad (3.2.4)$$

The charming part about (3.2.4) is that not a single aspect of it makes any sense. We know that $\mathbf{W}$ -a.s. $\omega \in \Omega$ does not even belong to $\mathcal{H}$, since Brownian paths are non-differentiable (in fact, they are Hölder continuous of any exponent less than $1/2$, but no better). Also, (3.2.4) tries to express the density of $\mathbf{W}$, but with respect to what measure? We have just stated that there is no “Lebesgue measure” on $\mathcal{H}$.

Despite these objections, Theorem 3.2.1 is a perfectly rigorous manifestation of the intuition that $\mathbf{W}$ is a standard Gaussian measure on $\mathcal{H}$. To reconcile this, it will turn out that a “standard $\mathcal{H}$ -Gaussian measure” *can exist*, but the catch is that it no longer “fits” in $\mathcal{H}$ (indeed, $\mathbf{W}$ is supported on $C([0, T]) \supsetneq \mathcal{H}$). Indeed, the fact that $\mathbf{W}$ is usually defined on $C([0, T])$ is somewhat of a red herring, and many of the deeper properties of Brownian motion (e.g., Schilder’s theorem in large deviations) are best understood via $\mathcal{H}$.

Abstract Wiener space.

More generally, let $\mathcal{H}$ be an infinite-dimensional Hilbert space and let us try to construct the standard Gaussian measure on $\mathcal{H}$. We tried earlier to use the sum $\sum_{k \in \mathbb{N}} \xi_k e_k$, but this does not converge in the norm of $\mathcal{H}$. To proceed, the idea is rather simple: we can just use another norm. Namely, if we can find a Banach space $\mathcal{B} \supseteq \mathcal{H}$ with corresponding norm $\|\cdot\|_{\mathcal{B}}$, such that the sum $\sum_{k \in \mathbb{N}} \xi_k e_k$ converges in $\|\cdot\|_{\mathcal{B}}$, then $\sum_{k \in \mathbb{N}} \xi_k e_k$ makes sense as a random element of $\mathcal{B}$, and we can take μ to be its law. Note

that μ is supported on $\mathcal{B}$ rather than on $\mathcal{H}$. It turns out that a suitable orthonormal basis $(e_k)_{k \in \mathbb{N}}$ and a norm $\|\cdot\|_{\mathcal{B}}$ can always be found to make this procedure work.

For Brownian motion, we take $\|\cdot\|_{\mathcal{B}}$ to be the supremum norm (i.e., the norm of $C([0, T])$) and the orthonormal basis can be chosen as a certain (integrated) wavelet basis. In fact, the abstract construction described above is implicit in Lévy's usual construction of Brownian motion.

It is also possible to flip this process around. Namely, suppose that μ is a Gaussian measure on a Banach space $\mathcal{B}$, which means that for any linear functionals $\ell_1, \dots, \ell_n \in \mathcal{B}^*$ and $X \sim \mu$, the vector $(\ell_1(X), \dots, \ell_n(X))$ is jointly Gaussian. Then, one can find a Hilbert space $\mathcal{H}$ associated to $(\mathcal{B}, \mu)$, which is called the Cameron–Martin space, and an appropriate analogue of the Cameron–Martin theorem (Theorem 3.2.1) holds. In fact, one has $\|x\|_{\mathcal{H}} := \sup\{|\ell(x)| \mid \ell \in \mathcal{B}^*, \|\ell\|_{L^2(\mu)} \leq 1\}$ and $\mathcal{H} := \{x \in \mathcal{B} \mid \|x\|_{\mathcal{H}} < \infty\}$. The triple $(\mathcal{B}, \mathcal{H}, \mu)$ is known as an **abstract Wiener space**.

3.2.2 Girsanov's Theorem

The Cameron–Martin theorem (Theorem 3.2.1) provides a formula for the density ratio of the laws of two diffusions that differ by a deterministic drift. Girsanov's theorem generalizes this to diffusions which differ by a random drift.

Why do we only consider a change of drift? The answer is that two diffusions

$$\begin{aligned} dX_t &= b_t^X dt + \sigma_t^X dB_t \\ dY_t &= b_t^Y dt + \sigma_t^Y dB_t \end{aligned}$$

must have mutually *singular* laws, unless their quadratic variations agree almost surely. Indeed, by Theorem 3.1.1, the laws of X and Y are concentrated on the events that the quadratic variation equals $\int_0^\cdot \sigma^X (\sigma^X)^\top$ or $\int_0^\cdot \sigma^Y (\sigma^Y)^\top$ respectively.

To understand the intuition behind Girsanov's theorem, consider the diffusion

$$dX_t = b_t dt + dB_t \iff X_t = \int_0^t b_s ds + B_t, \quad (3.2.5)$$

where $(b_t)_{t \geq 0}$ is an adapted process. If $(b_t)_{t \geq 0}$ is in fact deterministic, then the law of X_t is a Gaussian with mean $\int_0^t b_s ds$ and covariance tI_d , and is therefore fully explicit. In general, however, the law of X_t is not easily describable. Nevertheless, it is possible to understand the *joint law* of $(X_t)_{t \in [0, T]}$, which is a measure $\mathbf{P}$ on path space.

To see why this is the case, consider a discrete-time analogue of (3.2.5):

$$X_{n+1} := F(X_n) + \xi_n, \quad n = 0, 1, 2, \dots,$$

where $F : \mathbb{R}^d \rightarrow \mathbb{R}^d$ is a deterministic map and $(\xi_n)_{n \in \mathbb{N}}$ is a sequence of i.i.d. Gaussian variables. Again, due to the non-linear mapping F , the law of X_n is hard to describe exactly, yet once we *condition* on X_n , the law of X_{n+1} is an explicit Gaussian. This observation makes it straightforward to write down an explicit and simple expression for the joint law of $(X_0, X_1, \dots, X_N)$. See Exercise 3.3 for further elaboration on this discrete-time approach.

Returning to (3.2.5), we can think of it in the same way: namely, conditioned on the past, the conditional law of the diffusion in the next instant is a Gaussian with some mean and covariance. Moreover, by using the formula for the density ratio of two Gaussians, one can guess the formula

$$\frac{d\mathbf{P}}{d\mathbf{W}} = \exp\left(\int_0^T \langle b_s, dB_s \rangle - \frac{1}{2} \int_0^T \|b_s\|^2 ds\right). \quad (3.2.6)$$

Note that this exactly mirrors [Theorem 3.2.1](#), except that now we allow for adapted processes $(b_t)_{t \geq 0}$.

Let us now discuss how we would establish (3.2.6) carefully. Actually, we will proceed in the opposite order from our informal discussion: we will first *define* $\mathbf{P}$ via the formula (3.2.6), and then investigate the effect of the change of measure from $\mathbf{W}$ to $\mathbf{P}$ on our stochastic processes. This requires a change of perspective: instead of considering two processes $(B_t)_{t \in [0, T]}$ and $(X_t)_{t \in [0, T]}$, we instead think of a single process $(B_t)_{t \in [0, T]}$ defined on our canonical filtered space $(\Omega = C([0, T]), \mathcal{F}, (\mathcal{F}_t)_{t \in [0, T]})$. Recall that $(B_t)_{t \in [0, T]}$ is just the coordinate process $B_t(\omega) = \omega_t$. When we endow our space with the Wiener measure $\mathbf{W}$, then $(B_t)_{t \in [0, T]}$ becomes a standard Brownian motion. On the other hand, if we instead endow our space with the measure $\mathbf{P}$, we will show that $(B_t)_{t \in [0, T]}$ is, in some sense, a Brownian motion *with drift*.

To carry out our plan, the first step is to show that (3.2.6) defines a valid probability measure $\mathbf{P}$. In other words, we need the $\mathbf{W}$ -expectation of the right-hand side of (3.2.6) to equal 1. Actually, for any $t \in [0, T]$, let us write $\mathbf{W}_t$ for the restriction of $\mathbf{W}$ to $\mathcal{F}_t$ (and similarly write $\mathbf{P}_t$). In order for our putative $\mathbf{P}_t$ to be a probability measure for each $t \in [0, T]$, we would require

$$t \mapsto \frac{d\mathbf{P}_t}{d\mathbf{W}_t} \stackrel{?}{=} \exp\left(\int_0^t \langle b_s, dB_s \rangle - \frac{1}{2} \int_0^t \|b_s\|^2 ds\right)$$

to have constant $\mathbf{W}$ -expectation, equal to 1 for all $t \in [0, T]$. This would follow if we knew that this defined a $\mathbf{W}$ -martingale.

Assume that $\mathbb{E}^{\mathbf{W}} \int_0^T \|b_s\|^2 ds < \infty$. Then, the process $t \mapsto \int_0^t \langle b_s, dB_s \rangle$ is a $\mathbf{W}$ -martingale, and $t \mapsto \int_0^t \|b_s\|^2 ds$ is its quadratic variation. We will simply write $M := \int_0^\cdot \langle b, dB \rangle$ for the martingale and $[M, M]$ for its quadratic variation. Then,

$$\mathcal{E}(M) := \exp\left(M - \frac{1}{2} [M, M]\right) \quad (3.2.7)$$

is called the **exponential martingale** associated with M . Is it actually a martingale? Applying Itô's formula in the form (3.1.6),

$$d\mathcal{E}(M)_t = \mathcal{E}(M)_t \left(dM_t - \frac{1}{2} d[M, M]_t + \frac{1}{2} d[M, M]_t \right) = \mathcal{E}(M)_t dM_t,$$

so $\mathcal{E}(M)$ is a stochastic integral. From [Proposition 1.1.16](#), this tells us that $\mathcal{E}(M)$ is a continuous *local* $\mathbf{W}$ -martingale. In other words, it is possible for $\mathcal{E}(M)$ to *fail* to be a martingale if some integrability conditions are violated. For now, we will *assume* that $\mathcal{E}(M)$ is an honest³ martingale, treating this point as a technical issue, although later we will see that there is a clear understanding of what happens when this assumption fails.

Under this assumption, the measure $\mathbf{P}$ defined via (3.2.6) is a probability measure on path space. We now claim that under $\mathbf{P}$, the process $t \mapsto \tilde{B}_t := B_t - [B, M]_t = B_t - \int_0^t b_s ds$ is a standard Brownian motion. Actually, this is not too hard to check using (3.2.6) and characteristic functions; we leave it as [Exercise 3.2](#). It leads to the following theorem.

³We borrow the terminology from [Steele \(2001\)](#).

Theorem 3.2.8 (Girsanov). *Let $(B_t)_{t \in [0, T]}$ be a standard Brownian motion under the Wiener measure $\mathbf{W}$ and let $(b_t)_{t \in [0, T]}$ be a progressive process with $\mathbb{E}^{\mathbf{W}} \int_0^T \|b_s\|^2 ds < \infty$. Let $M_t := \int_0^t \langle b_s, dB_s \rangle$ for $t \in [0, T]$ and let $[M, M] = \int_0^t \|b_s\|^2 ds$ denote the quadratic variation. Define the exponential martingale*

$$\mathcal{E}(M) := \exp\left(M - \frac{1}{2} [M, M]\right).$$

Assume that $\mathcal{E}(M)$ is a $\mathbf{W}$ -martingale and define the measure $\mathbf{P}$ on path space via

$$\frac{d\mathbf{P}}{d\mathbf{W}} = \mathcal{E}(M)_T.$$

Then, under $\mathbf{P}$,

$$t \mapsto \tilde{B}_t := B_t - [B, M]_t = B_t - \int_0^t b_s ds \quad \text{is a standard Brownian motion.}$$

At present, Girsanov's theorem may seem rather abstract, and perhaps the best way to learn its meaning is to see it in action. We will put it to work in Section 4.4.

When is the exponential martingale an honest martingale?

Since $\mathcal{E}(M)$ is a non-negative local martingale, it is a *supermartingale*, i.e., we always have $\mathbb{E}^{\mathbf{W}} \mathcal{E}(M)_T \leq 1$. The only situation in which we encounter difficulties is when $\mathbb{E}^{\mathbf{W}} \mathcal{E}(M)_T < 1$, which would lead $\mathbf{P}$ defined via (3.2.6) to be a *sub-probability* measure. One might suspect that this is related to some probability mass “running off to ∞ ”, and indeed $1 - \mathbb{E}^{\mathbf{W}} \mathcal{E}(M)_T$ often corresponds to the probability that the diffusion has exploded by time T , in the sense discussed in Section 1.1.3. Therefore, the following criteria for $\mathcal{E}(M)$ to be a martingale are really criteria for non-explosion.

The standard sufficient condition for $\mathcal{E}(M)$ to be a martingale is **Novikov's condition**, which reads $\mathbb{E}^{\mathbf{W}} \exp(\frac{1}{2} [M, M]_T) < \infty$. An even weaker condition, known as **Kazamaki's condition**, requires only that $\sup_{t \in [0, T]} \mathbb{E}^{\mathbf{W}} \exp(\frac{1}{2} M_t) < \infty$. In principle, one of these conditions should be checked before applying Girsanov's theorem. However, if one is only interested in bounding a quantity such as the KL divergence or a Rényi divergence, one can often use the technique of localization, mentioned in Section 1.1.1, together with lower semicontinuity, to avoid these conditions altogether. In this book, we omit discussion of these technical arguments.

3.3 Doob's Transform

In this section, we will introduce a more sophisticated use of change of measure on path space, known as *Doob's transform*. Some applications include obtaining SDEs for processes conditioned on an endpoint and deriving the Föllmer process in the next section.

Suppose that $\mathbf{Q}$ is a reference measure on path space, describing the law of the SDE

$$dX_t = b_t(X_t) dt + \sigma_t(X_t) dB_t, \quad X_0 \sim \pi_0. \quad (3.3.1)$$

Let $P_{s,t}$ denote the transition operator from time s to time t (note that our reference process is time-inhomogeneous). The question we address in this section is the following: suppose that $\mathbf{P}$ is another probability measure on path space such that its Radon–Nikodym derivative w.r.t. $\mathbf{Q}$ only

depends on X_T , the path at time T :

$$\frac{d\mathbf{P}_T}{d\mathbf{Q}_T} = h_T(X_T).$$

Since $\mathbf{P}_T \ll \mathbf{Q}_T$, we know from the previous section that the process X under $\mathbf{P}_T$ corresponds to the original process under $\mathbf{Q}_T$ with a change of drift. Can we solve for this drift?

The Radon–Nikodym process $t \mapsto \frac{d\mathbf{P}_t}{d\mathbf{Q}_t}$ must be a martingale, and we make the ansatz that at time t , it only depends on the path at time t : $\frac{d\mathbf{P}_t}{d\mathbf{Q}_t} = h_t(X_t)$. Applying Itô's formula to $(h_t(X_t))_{t \in [0, T]}$,

$$dh_t(X_t) = (\partial_t h_t + \mathcal{L}_t h_t)(X_t) dt + \langle \nabla h_t(X_t), \sigma_t(X_t) dB_t \rangle, \quad (3.3.2)$$

and we deduce that this is a martingale provided

$$\partial_t h_t + \mathcal{L}_t h_t = 0,$$

where $\mathcal{L}_t f := \frac{1}{2} \langle \sigma_t \sigma_t^\top, \nabla^2 f \rangle + \langle b_t, \nabla f \rangle$ is the generator at time t . This is the *backward* heat equation (indeed, if X is a Brownian motion, then the equation reads $\partial_t h_t + \frac{1}{2} \Delta h_t = 0$), which makes sense since we have a terminal time condition for h_T . In the case when the process is time-homogeneous (i.e., the coefficients do not depend on t), setting $h_t^\leftarrow := h_{T-t}$, we see that $h^\leftarrow$ satisfies the forward heat equation $\partial_t h_t^\leftarrow = \mathcal{L} h_t^\leftarrow$, which has the solution $h_t^\leftarrow = P_{0,t} h_0^\leftarrow$. Switching back to h , we deduce that $h_t = P_{0, T-t} h_T$.

Now that we have a formula for $\frac{d\mathbf{P}_t}{d\mathbf{Q}_t}$, let us solve for the change of drift. On one hand, we know that if $\tilde{B} = B - [B, M]$ is a standard Brownian motion under $\mathbf{P}$, where $M := \int_0^\cdot \langle \Delta, dB \rangle$, then Girsanov's theorem (Theorem 3.2.8) yields

$$d\left(\log \frac{d\mathbf{P}_t}{d\mathbf{Q}_t}\right) = \langle \Delta_t, dB_t \rangle - \frac{1}{2} \|\Delta_t\|^2 dt.$$

On the other hand, Itô's formula and (3.3.2) yield

$$\begin{aligned} d(\log h_t(X_t)) &= \frac{1}{h_t(X_t)} dh_t(X_t) - \frac{1}{2 h_t(X_t)^2} d[h_t(X_t), h_t(X_t)]_t \\ &= \frac{1}{h_t(X_t)} \langle \nabla h_t(X_t), \sigma_t(X_t) dB_t \rangle - \frac{1}{2 h_t(X_t)^2} \|\sigma_t(X_t)^\top \nabla h_t(X_t)\|^2 dt \end{aligned}$$

and we quickly deduce that

$$\Delta_t = \sigma_t(X_t)^\top \nabla \log h_t(X_t).$$

Finally, we have obtained

$$\begin{aligned} dX_t &= (b_t(X_t) + \sigma_t(X_t) \Delta_t) dt + \sigma_t(X_t) d\tilde{B}_t \\ &= (b_t(X_t) + \sigma_t(X_t) \sigma_t(X_t)^\top \nabla \log h_t(X_t)) dt + \sigma_t(X_t) d\tilde{B}_t. \end{aligned}$$

This is our expression for the SDE under $\mathbf{P}$. We recall that in the time-homogeneous case, we actually have $h_t = P_{0, T-t} h_T$.

3.3.1 Conditioning on an Endpoint

As a first application, we will show how the Doob transform allows us to condition on $X_T = x_T$, for some fixed $x_T \in \mathbb{R}^d$. To see what this means, let $t < T$ and let η_t be a bounded $\mathcal{F}_t$ -measurable

random variable; let us try to compute $\mathbb{E}^{\mathbf{Q}}[\eta_t | X_T]$. By the definition of the conditional expectation, we wish to compute $\mathbb{E}^{\mathbf{Q}}[\eta_t f(X_T)]$ for any bounded measurable $f : \mathbb{R}^d \rightarrow \mathbb{R}$. Using the transition operator $P_{t,T}$ (which, in an abuse of notation, we identify with the transition kernel itself), and writing $\pi_t := \text{law}^{\mathbf{Q}}(X_t)$,

$$\begin{aligned} \mathbb{E}^{\mathbf{Q}}[\eta_t f(X_T)] &= \mathbb{E}^{\mathbf{Q}}[\eta_t \mathbb{E}^{\mathbf{Q}}[f(X_T) | X_t]] = \mathbb{E}^{\mathbf{Q}}\left[\eta_t \int f(x_T) P_{t,T}(X_t, dx_T)\right] \\ &= \mathbb{E}^{\mathbf{Q}}\left[\eta_t \int f(x_T) \frac{dP_{t,T}(X_t, \cdot)}{d\pi_T}(x_T) \pi_T(dx_T)\right] \\ &= \int f(x_T) \mathbb{E}^{\mathbf{Q}}\left[\eta_t \frac{dP_{t,T}(X_t, \cdot)}{d\pi_T}(x_T)\right] \pi_T(dx_T). \end{aligned}$$

We conclude that if $h_t^{x_T}(x) := \frac{dP_{t,T}(x, \cdot)}{d\pi_T}(x_T)$, then

$$\mathbb{E}^{\mathbf{Q}}[\eta_t | X_T = x_T] = \mathbb{E}^{\mathbf{Q}}[\eta_t h_t^{x_T}(X_t)].$$

Since this holds for every η_t , it says that if $\mathbf{P}$ denotes the measure $\mathbf{Q}$ conditioned on $X_T = x_T$, then

$$\frac{d\mathbf{P}_t}{d\mathbf{Q}_t} = h_t^{x_T}(X_t) \quad \text{for all } 0 \leq t < T.$$

We are now in the setting of Doob's transform. In particular, under $\mathbf{P}$,

$$dX_t = (b_t(X_t) + \sigma_t(X_t) \sigma_t(X_t)^\top \nabla \log h_t^{x_T}(X_t)) dt + \sigma_t(X_t) d\tilde{B}_t, \quad 0 \leq t < T. \quad (3.3.3)$$

Example 3.3.4 (Brownian bridge). Suppose that $B = X$ is a standard Brownian motion under $\mathbf{Q}$. Then, $P_{s,t}(x, \cdot) = \text{normal}(x, (t-s)I_d)$ and $\pi_t = P_{0,t}$. If we condition B to hit x_T at time T , then this process is known as a **Brownian bridge**.

We can calculate

$$h_t^{x_T}(x) \propto \exp\left(-\frac{\|x - x_T\|^2}{2(T-t)}\right) \implies \nabla \log h_t^{x_T}(x) = -\frac{x - x_T}{T-t}.$$

Hence, we arrive at the SDE representation for Brownian bridge,

$$dX_t = \frac{x_T - X_t}{T-t} dt + d\tilde{B}_t.$$

Note that the drift is singular as $t \nearrow T$. Of course, it must be, in order to drive the process to hit a single point x_T at time T .

3.3.2 Reversing the SDE

Next, we describe how to construct the *time reversal* of the SDE $(X_t)_{t \in [0, T]}$. Namely, suppose that $\pi_t := \text{law}(X_t)$ is the marginal law of X at time t . We will construct another SDE $X^\leftarrow$ such that $\text{law}(X_t^\leftarrow) = \pi_{T-t}$ for all $t \in [0, T]$. This construction will play an important role in the study of the proximal sampler in Chapter 8.

Perhaps the most straightforward approach is to start with the Fokker–Planck equation for X . Let X denote the general SDE (3.3.1), but for the sake of simplifying calculations we shall assume that

the diffusion matrix σ_t does not depend on the spatial variable, for all $t \in [0, T]$. Then, we have the Fokker–Planck equation

$$\partial_t \pi_t = \frac{1}{2} \langle \sigma_t \sigma_t^\top, \nabla^2 \pi_t \rangle - \operatorname{div}(\pi_t b_t).$$

Therefore, the time reversal $\pi_t^\leftarrow := \pi_{T-t}$ satisfies

$$\partial_t \pi_t^\leftarrow = -\frac{1}{2} \langle \sigma_{T-t} \sigma_{T-t}^\top, \nabla^2 \pi_t^\leftarrow \rangle + \operatorname{div}(\pi_t^\leftarrow b_{T-t}).$$

Next, we note that

$$\langle \sigma_{T-t} \sigma_{T-t}^\top, \nabla^2 \pi_t^\leftarrow \rangle = \operatorname{div}(\sigma_{T-t} \sigma_{T-t}^\top \nabla \pi_t^\leftarrow) = \operatorname{div}(\pi_t^\leftarrow (\sigma_{T-t} \sigma_{T-t}^\top \nabla \log \pi_t^\leftarrow)),$$

and hence we can write

$$\partial_t \pi_t^\leftarrow = \frac{1}{2} \langle \sigma_{T-t} \sigma_{T-t}^\top, \nabla^2 \pi_t^\leftarrow \rangle + \operatorname{div}(\pi_t^\leftarrow (b_{T-t} - \sigma_{T-t} \sigma_{T-t}^\top \nabla \log \pi_t^\leftarrow)). \quad (3.3.5)$$

We can therefore read off the SDE

$$dX_t^\leftarrow = \{-b_{T-t}(X_t^\leftarrow) + \sigma_{T-t} \sigma_{T-t}^\top \nabla \log \pi_t^\leftarrow(X_t^\leftarrow)\} dt + \sigma_{T-t} dB_t. \quad (3.3.6)$$

If we initialize this SDE with $X_0^\leftarrow \sim \pi_T$, then $X_t^\leftarrow \sim \pi_t^\leftarrow = \pi_{T-t}$ for all $t \in [0, T]$.

If we initialize this SDE at $X_0^\leftarrow = x_T$, does it then follow that $X_T^\leftarrow \sim \pi_{0|T}(\cdot | x_T)$, where $\pi_{0|T}$ denotes the conditional distribution of X_0 given X_T ? This would follow if we knew that $(X_0^\leftarrow, X_T^\leftarrow)$ has the same joint distribution as (X_T, X_0) , but this is not clear from the above derivation, which produced the process $X^\leftarrow$ by matching only the *marginal* laws. To see that this statement indeed holds, we will instead apply the conditioning argument from the previous section.

In the previous section, recalling that $\mathbf{P}$ is the measure $\mathbf{Q}$ conditioned on $X_T = x_T$, we know that $\operatorname{law}^\mathbf{P}(X_t) = \pi_{t|T}(\cdot | x_T)$. Therefore, from (3.3.3) and writing $\pi_{t|T} := \pi_{t|T}(\cdot | x_T)$ to lighten the notation, we deduce that

$$\partial_t \pi_{t|T} = \frac{1}{2} \langle \sigma_t \sigma_t^\top, \nabla^2 \pi_{t|T} \rangle - \operatorname{div}\left(\pi_{t|T} (b_t + \sigma_t \sigma_t^\top \nabla \log \frac{dP_{t,T}}{d\pi_T})\right).$$

The time reversal $\pi_{t|T}^\leftarrow := \pi_{T-t|T}$ satisfies

$$\partial_t \pi_{t|T}^\leftarrow = -\frac{1}{2} \langle \sigma_{T-t} \sigma_{T-t}^\top, \nabla^2 \pi_{t|T}^\leftarrow \rangle + \operatorname{div}\left(\pi_{t|T}^\leftarrow (b_{T-t} + \sigma_{T-t} \sigma_{T-t}^\top \nabla \log \frac{dP_{T-t,T}}{d\pi_T})\right). \quad (3.3.7)$$

As an application of the Bayes rule,

$$\frac{dP_{T-t,T}(x, \cdot)}{d\pi_T}(x_T) = \frac{\pi_{T|T-t}(x_T | x)}{\pi_T(x_T)} = \frac{\pi_{T-t|T}(x | x_T)}{\pi_{T-t}(x)}$$

and

$$\nabla \log \frac{\pi_{T-t|T}(\cdot | x_T)}{\pi_{T-t}} = \nabla \log \pi_{t|T}^\leftarrow - \nabla \log \pi_t^\leftarrow.$$

Substituting this into (3.3.7) and applying the logarithmic derivative trick,

$$\partial_t \pi_{t|T}^\leftarrow = \frac{1}{2} \langle \sigma_{T-t} \sigma_{T-t}^\top, \nabla^2 \pi_{t|T}^\leftarrow \rangle + \operatorname{div}(\pi_{t|T}^\leftarrow (b_{T-t} - \sigma_{T-t} \sigma_{T-t}^\top \nabla \log \pi_t^\leftarrow)). \quad (3.3.8)$$

Observe that this is the same Fokker–Planck equation as (3.3.5), except that it is now satisfied by

$t \mapsto \pi_{t|T}^\leftarrow$ rather than $t \mapsto \pi_t^\leftarrow$. Therefore, the SDE corresponding to (3.3.8) is the same SDE (3.3.6) above. Since $\pi_{0|T}^\leftarrow = \delta_{x_T}$, this confirms that initializing (3.3.6) with $X_0^\leftarrow = x_T$ yields $X_t^\leftarrow \sim \pi_{t|T}^\leftarrow$.

Note that *the time reversal of the SDE depends on the initial distribution*.

3.4 Föllmer Drift

As an application of the Doob transform, we now introduce a process attributed to Föllmer (1985). Under $\mathbf{Q}$, let $(X_t)_{t \in [0,1]}$ be a standard Brownian motion started at 0. Then, the law of X_1 is standard Gaussian, $\gamma := \text{normal}(0, I_d)$. Next, we take $\mathbf{P}$ to be a path measure under which $X_1 \sim \mu$, for some other probability measure μ with $\mu \ll \gamma$.

To achieve this, we can use the construction in Section 3.3, namely, we take $\frac{d\mathbf{P}_1}{d\mathbf{Q}_1} = h_1(X_1)$ where $h_1 = \frac{d\mu}{d\gamma}$. This yields $h_t = P_{1-t} \frac{d\mu}{d\gamma}$, and

$$dX_t = \nabla \log P_{1-t} \frac{d\mu}{d\gamma}(X_t) dt + d\tilde{B}_t, \quad X_0 = 0,$$

where $\tilde{B}$ is the $\mathbf{P}$ -Brownian motion. This process is known as the **Föllmer process**, and the added drift term is known as the **Föllmer drift**. The fundamental property enjoyed by this process (or more specifically, by $\mathbf{P}$) is that

$$\text{KL}(\mathbf{P} \parallel \mathbf{Q}) = \mathbb{E}^{\mathbf{P}} \log \frac{d\mu}{d\gamma}(X_1) = \mathbb{E}_\mu \log \frac{d\mu}{d\gamma} = \text{KL}(\mu \parallel \gamma). \quad (3.4.1)$$

On the other hand, if $\hat{\mathbf{P}}$ is any other path measure under which $X_1 \sim \mu$, then by the data-processing inequality (Theorem 1.5.6) we have $\text{KL}(\hat{\mathbf{P}} \parallel \mathbf{Q}) \geq \text{KL}(\mu \parallel \gamma)$. This reflects a certain *entropy optimality* property for the Föllmer process. Moreover, if $\hat{B}$ is a $\hat{\mathbf{P}}$ -Brownian motion and under $\hat{\mathbf{P}}$,

$$dX_t = \hat{b}_t(X_t) dt + d\hat{B}_t,$$

then by Girsanov's theorem (Theorem 3.2.8), this entropy optimality property becomes

$$\text{KL}(\mathbf{P} \parallel \mathbf{Q}) = \frac{1}{2} \mathbb{E}^{\mathbf{P}} \int_0^1 \left\| \nabla \log P_{1-t} \frac{d\mu}{d\gamma}(X_t) \right\|^2 dt \leq \frac{1}{2} \mathbb{E}^{\hat{\mathbf{P}}} \int_0^1 \|\hat{b}_t(X_t)\|^2 dt = \text{KL}(\hat{\mathbf{P}} \parallel \mathbf{Q}). \quad (3.4.2)$$

Thus, among all drifts that drive the process to satisfy $X_1 \sim \mu$, the Föllmer drift has minimal “energy”.

Moreover, by the chain rule for the KL divergence,

$$\text{KL}(\mathbf{P} \parallel \mathbf{Q}) = \text{KL}(\mu \parallel \gamma) + \int \text{KL}(\mathbf{P}^{X_1=x} \parallel \mathbf{Q}^{X_1=x}) \mu(dx),$$

where $\mathbf{P}^{X_1=x} = \text{law}_{\mathbf{P}}((X_t)_{0 \leq t < 1} \mid X_1 = x)$ and similarly for $\mathbf{Q}^{X_1=x}$. The optimality property (3.4.1) for the Föllmer process entails that the second term above vanishes, which means that under $\mathbf{P}$, the conditional law of the path given $X_1 = x$ is the same as the corresponding conditional law under $\mathbf{Q}$. The latter corresponds to the Brownian bridge (see Example 3.3.4).

3.4.1 Application to Functional Inequalities

The use of the Föllmer drift as a potent tool for establishing functional inequalities was perhaps pioneered by Lehec (2013), although he attributes the idea earlier, e.g., to Borell (2000). Here, we demonstrate its power to establish the Gaussian log-Sobolev and T_2 inequalities, and refer to Lehec (2013) and subsequent literature for further applications.

Transport inequality.

Under $\mathbf{P}$, we know that $X_1 \sim \mu$ and $\tilde{B}_1 \sim \gamma$. Hence,

$$\begin{aligned} W_2^2(\mu, \gamma) &\leq \mathbb{E}^{\mathbf{P}}[\|X_1 - \tilde{B}_1\|^2] = \mathbb{E}^{\mathbf{P}}\left[\left\|\int_0^1 \nabla \log P_{1-t} \frac{d\mu}{d\gamma}(X_t) dt\right\|^2\right] \\ &\leq \mathbb{E}^{\mathbf{P}} \int_0^1 \left\|\nabla \log P_{1-t} \frac{d\mu}{d\gamma}(X_t)\right\|^2 dt = 2 \text{KL}(\mu \parallel \gamma), \end{aligned}$$

where we used (3.4.1) and (3.4.2) in the last line.

Log-Sobolev inequality.

Let $h_t := P_{1-t} \frac{d\mu}{d\gamma}$ and recall from the construction of the Doob transform that $(h_t(X_t))_{t \in [0,1]}$ is a $\mathbf{Q}$ -martingale. We claim that $(\nabla \log h_t(X_t))_{t \in [0,1]}$ is a $\mathbf{P}$ -martingale. To prove this, let $s \leq t$ and let $A_s \in \mathcal{F}_s$. Then,

$$\begin{aligned} \mathbb{E}^{\mathbf{P}}[\nabla \log h_t(X_t) \mathbb{1}_{A_s}] &= \mathbb{E}^{\mathbf{P}}\left[\frac{\nabla h_t(X_t)}{h_t(X_t)} \mathbb{1}_{A_s}\right] = \mathbb{E}^{\mathbf{Q}}[\nabla h_t(X_t) \mathbb{1}_{A_s}] = \mathbb{E}^{\mathbf{Q}}[P_{t-s} \nabla h_t(X_s) \mathbb{1}_{A_s}] \\ &= \mathbb{E}^{\mathbf{Q}}[\nabla P_{t-s} h_t(X_s) \mathbb{1}_{A_s}] = \mathbb{E}^{\mathbf{Q}}[\nabla h_s(X_s) \mathbb{1}_{A_s}] = \mathbb{E}^{\mathbf{P}}[\nabla \log h_s(X_s) \mathbb{1}_{A_s}]. \end{aligned}$$

This yields

$$\mathbb{E}^{\mathbf{P}}[\nabla \log h_t(X_t) \mid \mathcal{F}_s] = \nabla \log h_s(X_s).$$

In particular, $t \mapsto \|\nabla \log h_t(X_t)\|^2$ is a $\mathbf{P}$ -submartingale.

Now, using the equality for the KL divergence and the submartingale property,

$$\text{KL}(\mu \parallel \gamma) = \frac{1}{2} \mathbb{E}^{\mathbf{P}} \int_0^1 \left\|\nabla \log P_{1-t} \frac{d\mu}{d\gamma}(X_t)\right\|^2 dt \leq \frac{1}{2} \mathbb{E}^{\mathbf{P}}\left[\left\|\nabla \log \frac{d\mu}{d\gamma}(X_1)\right\|^2\right] = \frac{1}{2} \text{Fl}(\mu \parallel \gamma).$$

3.4.2 Connection to Stochastic Localization

In this section, we relate the Föllmer process to Eldan's **stochastic localization** scheme (introduced in Eldan (2013)), which by now has solidified its status as a core tool in high-dimensional probability. Although we do not have space in this book to describe its many applications, we present some basic ideas here to help the reader understand the connections with the extant literature; see El Alaoui et al. (2022); Klartag and Putterman (2023) for more details.

Stochastic localization is a method of understanding a probability measure μ by decomposing it into simpler parts, with the goal of, e.g., establishing functional inequalities or other useful properties for μ . It was inspired by earlier work on deterministic localization schemes (Kannan et al., 1995), but instead seeks to produce a *random* measure-valued process $(p_t)_{t \geq 0}$. This process is such that $p_0 = \mu$, $p_\infty = \delta_X$ for some random variable X , and $(p_t)_{t \geq 0}$ is a martingale. The last property implies that $X \sim \mu$, and indeed, $\mathbb{E} p_t = \mu$ for all $t \geq 0$. This is the decomposition of μ into “simpler parts” as alluded to earlier.

How might we build such a process? We motivate the process via a Bayesian interpretation. Consider the process $\theta_t := tX + B_t$, where as usual $(B_t)_{t \geq 0}$ is a Brownian motion, independent of X . At time $t \approx 0$, the Brownian motion dominates, so θ_t contains almost no information about X . At time $t \rightarrow \infty$, the linear term tX dominates and θ_t contains nearly perfect information about X . Therefore, if we set p_t to be the conditional law of X given the observation θ_t , then we expect $p_0 = \mu$ and $p_\infty = \delta_X$. One can check that $(p_t)_{t \geq 0}$ is indeed a martingale.

We can relate this process to the usual heat flow via rescaling: let $\check{\theta}_t := \theta_t/t$, so that $\check{\theta}_t = X + B_t/t$. The time inversion property of Brownian motion implies that $(\check{B}_t)_{t \geq 0}$ is also a Brownian motion, where $\check{B}_{1/t} := B_t/t$. Hence, $\check{\theta}_t = X + \check{B}_{1/t}$ is the output of the heat flow, started at X , after time $1/t$ (note the time inversion). If we let $\rho_{0,s}$ denote the joint distribution of the heat flow, started at μ , at times 0 and s , and denote conditional distributions accordingly, we can write $p_t = \text{law}(X \mid \check{\theta}_t) = \rho_{0|t^{-1}}(\cdot \mid \theta_t/t)$. Thus, we have the explicit expression

$$p_t(dx) \propto \exp\left(-\frac{t \|x - \theta_t/t\|^2}{2}\right) \mu(dx).$$

We now want to take logarithms in this expression and apply Itô's formula to derive a stochastic evolution equation. Before doing so, note that the equation $\theta_t = tX + B_t$, which seems to give $d\theta_t = X dt + dB_t$, does not express θ as a Markov process (since X is not adapted to the filtration at time $t < \infty$). However, one can replace this equation by the equivalent Markov evolution $d\theta_t = \mathbb{E}[X \mid \mathcal{F}_t] dt + dW_t = a_t dt + dW_t$, where we set $a_t := \int x p_t(dx)$ and $(W_t)_{t \geq 0}$ is another Brownian motion; see [Klartag and Putterman \(2023\)](#). Then, a calculation starting from

$$d \log p_t(x) = -d\left(\frac{t \|x - \theta_t/t\|^2}{2}\right) - d \log \int \exp\left(-\frac{t \|y - \theta_t/t\|^2}{2}\right) \mu(dy)$$

eventually yields ([Exercise 3.6](#))

$$dp_t(x) = p_t(x) \langle x - a_t, dW_t \rangle, \quad (3.4.3)$$

which was the form in which stochastic localization was originally introduced.

To see the connection with the Föllmer process $(F_t)_{t \in [0,1]}$ with $F_1 \sim \mu$, recall that given F_1 , the law of $(F_t)_{0 \leq t \leq 1}$ is a Brownian bridge. In other words, $F_t = tF_1 + \mathbb{B}B_t$, where $(\mathbb{B}B_t)_{t \in [0,1]}$ is the Brownian bridge process starting and ending at 0. We can identify $F_1 = X$, in which case this expression for F nearly resembles the expression for the tilt process θ in stochastic localization. To make this precise, we claim that the Brownian bridge can be constructed as $\mathbb{B}B_t = (1-t) B_{t/(1-t)}$. With this identification, $F_t = (1-t) \theta_{t/(1-t)}$, i.e., the Föllmer process is a rescaled time compression of the tilt process θ from $\mathbb{R}_+$ to the interval $[0, 1]$; see [Exercise 3.7](#) for details.

This discussion also shows that these concepts are related to the idea of running the heat flow *backward* in time (e.g., we consider the conditional distribution $\rho_{0|t^{-1}}$ above). These ideas will reappear in [Chapter 8](#) as the proximal sampler.

3.5 Schrödinger Bridge

In this section, we consider a generalization of the Föllmer process. The setup arises from a hot gas *Gedankenexperiment* due to Schrödinger. Let μ and ν be two probability measures over $\mathbb{R}^d$, representing the observed distribution of a cloud of particles at times 0 and 1, respectively. In the absence of the observation of ν , we may have modelled the evolution of the gas particles as a scaled Brownian motion: $X_t = X_0 + \sqrt{\varepsilon} B_t$ for $t \in [0, 1]$, where $X_0 \sim \mu$ and $(B_t)_{t \in [0,1]}$ are independent, and $\varepsilon > 0$ represents the noise level of the process. However, if the observed distribution ν differs from the law $\mu * \text{normal}(0, \varepsilon I_d)$ of X_1 in our model, what then is our best guess for the law of the trajectory $(X_t)_{t \in [0,1]}$? The law of the trajectory is said to *bridge* the distributions μ and ν .

Schrödinger formulated this as a KL minimization problem:

$$\underset{\mathbf{P} \in \mathcal{P}(C([0,1]))}{\text{minimize}} \quad \text{KL}(\mathbf{P} \parallel \mathbf{W}^{\mu, \varepsilon}) \quad \text{such that} \quad (X_0)_\# \mathbf{P} = \mu, \quad (X_1)_\# \mathbf{P} = \nu.$$

Here, the minimization takes place over the set of path measures (probability measures over $C([0, 1])$), and $\mathbf{W}^{\mu, \varepsilon}$ denotes the path measure corresponding to Brownian motion, rescaled by $\sqrt{\varepsilon}$ and started at μ . The choice of KL divergence as our criterion can be motivated by large deviations theory.

We can solve this problem as follows. If we condition on the endpoints X_0 and X_1 and apply the KL chain rule, we end up with

$$\text{KL}(\mathbf{P} \parallel \mathbf{W}^{\mu, \varepsilon}) = \text{KL}(\text{law}_{\mathbf{P}}(X_0, X_1) \parallel \text{law}_{\mathbf{W}^{\mu, \varepsilon}}(X_0, X_1)) + \mathbb{E}^{\mathbf{P}} \text{KL}(\mathbf{P}^{X_0, X_1} \parallel \mathbf{W}^{\mu, \varepsilon|X_0, X_1}),$$

where $\mathbf{P}^{X_0, X_1}$, $\mathbf{W}^{\mu, \varepsilon|X_0, X_1}$ denote the path measures $\mathbf{P}$, $\mathbf{W}^{\mu, \varepsilon}$ conditioned on (X_0, X_1) respectively. Also, $\mathbf{W}^{\mu, \varepsilon|X_0, X_1}$ is a Brownian bridge (rescaled by $\sqrt{\varepsilon}$), and the second term above can be made zero by setting $\mathbf{P}^{X_0, X_1} = \mathbf{W}^{\mu, \varepsilon|X_0, X_1}$.

Thus far, the development has closely mirrored our discussion of the Föllmer process, which is the special case of the Schrödinger bridge when $\mu = \delta_0$. The interesting new features of this more general setting arise, however, when we consider minimizing the first term in the KL chain rule, which is a minimization over joint distributions for (X_0, X_1) with $X_0 \sim \mu$ and $X_1 \sim \nu$, i.e., a *coupling* of μ and ν . In the Föllmer case, this minimization problem was trivial, essentially because the space of couplings of a Dirac measure δ_0 and any other measure ν is also trivial (consisting solely of $\delta_0 \otimes \nu$). Our goal now is to understand this minimization problem when the space of couplings is non-trivial.

3.5.1 Entropically Regularized Optimal Transport

Our first step is to note that for $\eta := \text{law}_{\mathbf{W}^{\mu, \varepsilon}}(X_0, X_1)$,

$$\eta(dx, dy) \propto \mu(dx) \exp\left(-\frac{\|y - x\|^2}{2\varepsilon}\right) dy.$$

Therefore, we can explicitly write, for $\gamma := \text{law}_{\mathbf{P}}(X_0, X_1)$,

$$\begin{aligned} \text{KL}(\gamma \parallel \eta) &= \frac{1}{2\varepsilon} \int \|x - y\|^2 \gamma(dx, dy) + \int \log \frac{\gamma(x, y)}{\mu(x)} \gamma(dx, dy) + \text{const.} \\ &= \frac{1}{2\varepsilon} \int \|x - y\|^2 \gamma(dx, dy) + \text{KL}(\gamma \parallel \mu \otimes \nu) + \text{const.} \end{aligned}$$

where we used the fact that $\int \log \nu(y) \gamma(dx, dy) = \int \log \nu d\nu$ does not depend on γ , allowing us to absorb it into the constant term. Hence, the problem of finding the optimal coupling γ between μ and ν for the Schrödinger bridge problem is an entropically regularized variant of the optimal transport problem from Section 1.3. More generally, we have:

Definition 3.5.1. Let $\mathcal{X}$ and $\mathcal{Y}$ be complete separable metric spaces, let $\varepsilon > 0$, and let $c : \mathcal{X} \times \mathcal{Y} \rightarrow [0, \infty]$ be a cost function. The **entropically regularized optimal transport cost** from $\mu \in \mathcal{P}(\mathcal{X})$ to $\nu \in \mathcal{P}(\mathcal{Y})$ with cost c is

$$\mathcal{T}_{c, \varepsilon}(\mu, \nu) := \inf_{\gamma \in \mathcal{C}(\mu, \nu)} \left[\int c(x, y) \gamma(dx, dy) + \varepsilon \text{KL}(\gamma \parallel \mu \otimes \nu) \right].$$

One can show that if c is lower semicontinuous, then there always exists a unique entropic optimal transport plan. Note that unlike [Theorem 1.3.8](#), which required μ to have a density in order for uniqueness of the optimal transport plan with quadratic cost, here uniqueness always holds as a consequence of the strict convexity of the KL divergence.

To summarize our observations thus far, we have argued that the solution to the Schrödinger bridge problem is to first draw (X_0, X_1) from the entropic optimal transport plan between μ and ν with cost $c(x, y) = \frac{1}{2} \|x - y\|^2$, and then to join X_0 and X_1 by a Brownian bridge (rescaled by $\sqrt{\varepsilon}$). Although in principle this completes the description of the Schrödinger bridge, we can go further by characterizing the entropic optimal transport plan via a duality principle.

Duality works here similarly as Kantorovich duality did for unregularized optimal transport, and we simply quote the main theorem here.

Theorem 3.5.2 (Duality for entropic optimal transport). *There exist maximizers $f_\varepsilon, g_\varepsilon$ to the dual problem*

$$\sup_{(f, g) \in L^1(\mu) \times L^1(\nu)} \left\{ \int f \, d\mu + \int g \, d\nu - \varepsilon \iint \exp\left(\frac{f \oplus g - c}{\varepsilon}\right) d(\mu \otimes \nu) + \varepsilon \right\}$$

which are unique up to adding a constant to f_ε and subtracting that same constant from g_ε . The optimal value of the dual problem equals the entropic optimal transport cost from μ to ν , and the entropic optimal transport plan γ_ε is of the form

$$\gamma_\varepsilon(dx, dy) = \exp\left(\frac{f_\varepsilon(x) + g_\varepsilon(y) - c(x, y)}{\varepsilon}\right) \mu(dx) \nu(dy). \quad (3.5.3)$$

The expression (3.5.3) characterizes the optimal solution in the following sense. If γ_ε is any coupling of μ and ν of the form (3.5.3) for some $f_\varepsilon, g_\varepsilon$, then γ_ε is the entropic optimal transport plan, and $(f_\varepsilon, g_\varepsilon)$ is a pair of optimal dual potentials.

Note that compared to the dual problem for Kantorovich duality in (1.3.7), we have replaced the “hard” constraint of $f \oplus g \leq c$ with the “soft” constraint of adding the Lagrangian penalty term $\iint \exp((f \oplus g - c)/\varepsilon) d(\mu \otimes \nu)$ into the objective.

Let us now specialize to the case of quadratic cost, $c(x, y) = \frac{1}{2} \|x - y\|^2$. In this case, it is natural to work with $\varphi_\varepsilon := \frac{1}{2} \|\cdot\|^2 - f_\varepsilon$ and $\psi_\varepsilon := \frac{1}{2} \|\cdot\|^2 - g_\varepsilon$, since then (provided that we fix a normalization for the potentials, e.g., $\int \varphi_\varepsilon \, d\mu = \int \psi_\varepsilon \, d\nu$) we have the convergence $\varphi_\varepsilon \rightarrow \varphi$ and $\psi_\varepsilon \rightarrow \psi$ of the entropic potentials to their unregularized counterparts as $\varepsilon \searrow 0$ (see Nutz and Wiesel, 2022). The condition that γ_ε has marginals μ and ν yields the coupled equations

$$\begin{aligned} \varphi_\varepsilon(x) &= \varepsilon \log \int \exp\left(\frac{\langle x, y \rangle - \psi_\varepsilon(y)}{\varepsilon}\right) \nu(dy), \\ \psi_\varepsilon(y) &= \varepsilon \log \int \exp\left(\frac{\langle x, y \rangle - \varphi_\varepsilon(x)}{\varepsilon}\right) \mu(dx). \end{aligned} \quad (3.5.4)$$

From these expressions, one can prove the following lemma (see Exercise 3.10).

Lemma 3.5.5. *The following relations hold:*

$$\nabla \varphi_\varepsilon(x) = \mathbb{E}_{\gamma_\varepsilon}[Y \mid X = x], \quad \nabla \psi_\varepsilon(y) = \mathbb{E}_{\gamma_\varepsilon}[X \mid Y = y].$$

Also,

$$\nabla^2 \varphi_\varepsilon(x) = \frac{1}{\varepsilon} \text{cov}_{\gamma_\varepsilon}(Y \mid X = x), \quad \nabla^2 \psi_\varepsilon(y) = \frac{1}{\varepsilon} \text{cov}_{\gamma_\varepsilon}(X \mid Y = y).$$

Since covariance matrices are always positive semidefinite, these expressions witness the convexity of φ_ε and ψ_ε , and provide another explanation for Brenier's theorem (Theorem 1.3.8).

3.5.2 Caffarelli's Contraction Theorem

We now provide an application of the theory of entropic optimal transport to functional inequalities. We will establish bounds on the Hessians of the entropic potentials which, as $\varepsilon \searrow 0$, furnish bounds on the Brenier potential for the unregularized optimal transport problem. This will yield a proof of Caffarelli's contraction theorem.

Theorem 3.5.6 (Chewi and Pooladian (2023)). *Suppose that μ is β -log-smooth and that ν is α -strongly log-concave, i.e., $\nabla^2 \log(1/\mu) \leq \beta I_d$ and $\nabla^2 \log(1/\nu) \geq \alpha I_d$. Then, the entropic Brenier potential φ_ε from μ to ν satisfies*

$$\nabla^2 \varphi_\varepsilon \leq \frac{1}{2} (\sqrt{4\beta/\alpha + \varepsilon^2 \beta^2} - \varepsilon \beta) I_d.$$

Letting $\varepsilon \searrow 0$, one readily obtains (see Chewi and Pooladian, 2023):

Corollary 3.5.7 (Caffarelli's contraction theorem, Caffarelli (2000)). *Suppose that μ is β -log-smooth and that ν is α -strongly log-concave, i.e., $\nabla^2 \log(1/\mu) \leq \beta I_d$ and $\nabla^2 \log(1/\nu) \geq \alpha I_d$. Then, the Brenier map $\nabla \varphi$ from μ to ν is $\sqrt{\beta/\alpha}$ -Lipschitz.*

The proof of Theorem 3.5.6 will exploit the representation of the Hessians of the entropic Brenier potentials as covariance matrices (Lemma 3.5.5), together with a pair of covariance inequalities.

Theorem 3.5.8 (Cramér–Rao inequality). *Let $\pi \propto \exp(-V)$ be a probability measure over $\mathbb{R}^d$. For any well-behaved function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, it holds that*

$$\text{var}_\pi f \geq \langle \mathbb{E}_\pi \nabla f, (\mathbb{E}_\pi \nabla^2 V)^{-1} \mathbb{E}_\pi \nabla f \rangle.$$

Proof Integration by parts and $\mathbb{E}_\pi \nabla V = 0$ yield

$$\mathbb{E}_\pi \nabla f = \int \nabla f \, d\pi = \int (f \nabla \log \frac{1}{\pi}) \, d\pi = \mathbb{E}_\pi [(f - \mathbb{E}_\pi f) \nabla V].$$

Therefore,

$$\begin{aligned} \langle \mathbb{E}_\pi \nabla f, (\mathbb{E}_\pi \nabla^2 V)^{-1} \mathbb{E}_\pi \nabla f \rangle &= \mathbb{E}_\pi [(f - \mathbb{E}_\pi f) \langle \nabla V, (\mathbb{E}_\pi \nabla^2 V)^{-1} \mathbb{E}_\pi \nabla f \rangle] \\ &\leq \sqrt{(\text{var}_\pi f) \mathbb{E}_\pi \langle \mathbb{E}_\pi \nabla f, (\mathbb{E}_\pi \nabla^2 V)^{-1} (\nabla V)^{\otimes 2} (\mathbb{E}_\pi \nabla^2 V)^{-1} \mathbb{E}_\pi \nabla f \rangle}. \end{aligned}$$

Another integration by parts shows that $\mathbb{E}_\pi [(\nabla V)^{\otimes 2}] = \mathbb{E}_\pi \nabla^2 V$, so the result follows by rearranging the above expression. $\square$

Corollary 3.5.9 (Covariance bounds). *Let $\pi \propto \exp(-V)$ be a probability measure over $\mathbb{R}^d$ and let cov_π denote its covariance matrix. Then,*

$$(\mathbb{E}_\pi \nabla^2 V)^{-1} \leq \text{cov}_\pi \leq \mathbb{E}_\pi [(\nabla^2 V)^{-1}],$$

where the second bound holds if π is strictly log-concave.

Proof The lower and upper bounds follow respectively from the Cramér–Rao inequality ([Theorem 3.5.8](#)) and the Brascamp–Lieb inequality ([Theorem 2.2.9](#)) by taking test functions $f = \langle e, \cdot \rangle$ for unit vectors $e \in \mathbb{R}^d$. $\square$

Proof of Theorem 3.5.6 In the proof, we write $\mu \propto \exp(-V)$ and $\nu \propto \exp(-W)$. For any point $x \in \mathbb{R}^d$, from [Lemma 3.5.5](#) and the upper bound in [Corollary 3.5.9](#), we obtain

$$\nabla^2 \varphi_\varepsilon(x) = \varepsilon^{-1} \text{cov}_{\gamma_\varepsilon^{Y|X=x}} \leq \varepsilon^{-1} \mathbb{E}_{\gamma_\varepsilon^{Y|X=x}} [(\varepsilon^{-1} \nabla^2 \psi_\varepsilon + \nabla^2 W)^{-1}] \leq \mathbb{E}_{\gamma_\varepsilon^{Y|X=x}} [(\nabla^2 \psi_\varepsilon + \varepsilon \alpha I)^{-1}].$$

For any $y \in \mathbb{R}^d$, from [Lemma 3.5.5](#) and the lower bound in [Corollary 3.5.9](#),

$$\nabla^2 \psi_\varepsilon(y) = \varepsilon^{-1} \text{cov}_{\gamma_\varepsilon^{X|Y=y}} \geq \varepsilon^{-1} (\mathbb{E}_{\gamma_\varepsilon^{X|Y=y}} [\varepsilon^{-1} \nabla^2 \varphi_\varepsilon + \nabla^2 V])^{-1} \geq (\mathbb{E}_{\gamma_\varepsilon^{X|Y=y}} [\nabla^2 \varphi_\varepsilon + \varepsilon \beta I])^{-1}.$$

Let $L_\varepsilon := \sup_{x \in \mathbb{R}^d} \lambda_{\max}(\nabla^2 \varphi_\varepsilon(x))$. From the two inequalities above, we can conclude that

$$\lambda_{\max}(\nabla^2 \varphi_\varepsilon(x)) \leq ((L_\varepsilon + \varepsilon \beta)^{-1} + \varepsilon \alpha)^{-1}$$

and hence

$$L_\varepsilon \leq ((L_\varepsilon + \varepsilon \beta)^{-1} + \varepsilon \alpha)^{-1}.$$

Solving the quadratic inequality yields the upper bound on L_ε in the theorem. $\square$

As discussed in [Section 2.3](#), if we apply Caffarelli’s contraction theorem taking μ as the standard Gaussian measure (so $\beta = 1$), we deduce that the optimal transport map from the standard Gaussian to any α -strongly log-concave measure is $\alpha^{-1/2}$ -Lipschitz. Together with the preservation of functional inequalities under Lipschitz mappings ([Proposition 2.3.3](#)), it allows us to transfer functional inequalities satisfied by the standard Gaussian measure to all strongly log-concave measures. In this way, Caffarelli’s contraction theorem is a “universal blueprint” for proving such inequalities. This point was already made in Caffarelli’s original paper ([Caffarelli, 2000](#)).

For example, one can use it to transfer the functional inequalities established for the standard Gaussian in [Section 3.4](#) to strongly log-concave measures. As another example, one can prove the Gaussian isoperimetric inequality in [Theorem 2.4.29](#) by first establishing it for Gaussians and appealing to Caffarelli contraction. Another approach to constructing Lipschitz transport maps is explored in [Exercise 3.8](#).

Bibliographical Notes

The discussion of quadratic variation and Girsanov’s theorem is heavily inspired by the treatment in [Le Gall \(2016\)](#). The study of functions of bounded variation and the relationship with absolute continuity and total variation can be found in any standard graduate text on real analysis, e.g., [Folland \(1999\)](#). See [Steele \(2001, Proposition 8.6\)](#) for a simple case of [Theorem 3.1.5](#).

The Cameron–Martin theorem is named after [Cameron and Martin \(1944\)](#), and abstract Wiener

spaces were introduced in [Gross \(1967\)](#). This is the starting point of calculus on path space, which is usually called the **Malliavin calculus**. To illustrate, suppose we have a functional $F(B)$ of the Brownian path B , and we ask how $F(B)$ changes under infinitesimal perturbations $B \mapsto B + h$. The key point is that one should restrict to directions h which belong to the Cameron–Martin space, which leads to the concept of the Malliavin derivative.

Girsanov’s theorem is also used heavily in mathematical finance, and for this purpose the book [Steele \(2001\)](#) is warmly recommended.

The proof of [Exercise 3.4](#) was first given in [Arnaudon et al. \(2006\)](#). An extension to potentials which are only assumed to be strongly convex at infinity is presented in [Lu and Wang \(2025\)](#) via reflection coupling. See [Altschuler and Chewi \(2024b\)](#) for further discussion.

The time reversal of SDEs was studied in [Haussmann and Pardoux \(1986\)](#) and formulated rigorously under a finite entropy condition in [Föllmer \(1985\)](#). See [Cattiaux et al. \(2023\)](#) for an extended discussion. The proof of [Exercise 3.5](#) is taken from [Conforti et al. \(2025\)](#).

Stochastic localization was introduced by R. Eldan in [Eldan \(2013\)](#), and he and other researchers have subsequently used the Föllmer process and stochastic localization to great effect for various applications; see, e.g., [Eldan \(2015, 2016\)](#); [Eldan et al. \(2017\)](#); [Lee and Vempala \(2017\)](#); [Eldan \(2018\)](#); [Eldan and Lee \(2018\)](#); [Klartag \(2018\)](#); [Lee and Vempala \(2018\)](#); [Eldan \(2020\)](#); [Eldan et al. \(2020, 2022\)](#); [Eldan and Shamir \(2022\)](#); [Eldan \(2023\)](#); [Klartag \(2023\)](#); [Klartag and Putterman \(2023\)](#); [Klartag and Lehec \(2026\)](#). Algorithmic applications were developed in [El Alaoui et al. \(2022\)](#); [Montanari \(2023\)](#) and are closely related to the proximal sampler in Chapter 8. These ideas can be placed in the more general framework of localization schemes ([Chen and Eldan, 2022](#)), and can be interpreted in terms of renormalization and the Polchinski flow ([Bauerschmidt et al., 2024](#)). See also [Kook and Vempala \(2025b\)](#); [Shi et al. \(2025\)](#) for recent surveys.

The Kim–Milman map from [Exercise 3.8](#) was introduced in [Kim and Milman \(2012\)](#). It has been used to establish functional inequalities in subsequent works, e.g., [Klartag and Putterman \(2023\)](#); [Mikulincer and Shenfeld \(2023\)](#); [Brigati and Pedrotti \(2024\)](#); [Chewi et al. \(2024\)](#); [Fathi et al. \(2024\)](#). See [Exercise 3.9](#) and the closely related approach of the Brownian transport map ([Mikulincer and Shenfeld, 2024](#)). In [Exercise 3.9](#), it is shown that a log-concave distribution over a set of diameter R has log-Sobolev constant $O(R^2)$; if in addition the distribution is assumed to be isotropic, then the estimate improves to $O(R)$, which is tight ([Lee and Vempala, 2018](#)).

See [Chewi \(2024, Chapter 4\)](#) for background on entropic optimal transport, and [Léonard \(2012\)](#); [Chen et al. \(2021b\)](#) for introductions to the Schrödinger bridge problem. In recent years, entropic optimal transport has been popularized by [Cuturi \(2013\)](#) as a fast computational approach to optimal transport. An entropic optimal transport approach to Caffarelli’s contraction theorem was first given in [Fathi et al. \(2020\)](#), and the proof in Section 3.5.2 was generalized in [Gozlan and Sylvestre \(2025\)](#). Entropic optimal transport was also used to prove other functional inequalities in [Ledoux \(2018\)](#); [Gentil et al. \(2020\)](#). In the Gaussian case, closed-form expressions were computed for the entropic optimal transport cost in [Janati et al. \(2020\)](#); [Mallasto et al. \(2022\)](#), generalizing [Exercise 3.11](#), and for the Schrödinger bridge in [Bunne et al. \(2023\)](#).

Exercises

Change of Measure in Path Space

▷ Exercise 3.1 (Sobolev embedding)

Let $h \in \mathcal{H}$, where $\mathcal{H}$ is the Cameron–Martin space defined in (3.2.2). Prove that h is $1/2$ -Hölder continuous. *Hint:* Use the fundamental theorem of calculus and Cauchy–Schwarz.

This is quite suggestive, since Brownian paths are s -Hölder continuous for every $s < 1/2$.

▷ Exercise 3.2 (Proof of Girsanov’s theorem)

Let $0 = t_0 < t_1 < \dots < t_n \leq T$ and $\theta_1, \dots, \theta_n \in \mathbb{R}^d$. Let b , B , and $\tilde{B}$ be as in Theorem 3.2.8. Using the formula for the Radon–Nikodym derivative of $\mathbf{P}$ w.r.t. $\mathbf{W}$, compute $\mathbb{E}^{\mathbf{P}} \exp(\mathbf{i} \sum_{i=1}^n \langle \theta_i, \tilde{B}_{t_i} - \tilde{B}_{t_{i-1}} \rangle)$, where $\mathbf{i} := \sqrt{-1}$, and deduce that $\tilde{B}$ is a $\mathbf{P}$ -Brownian motion.

▷ Exercise 3.3 (Girsanov’s theorem in discrete time)

To gain intuition for Girsanov’s theorem, consider the following discrete approximations to two SDEs:

$$\begin{aligned} X_{n+1} &:= X_n - h b_n^X(X_n) + \sqrt{h} \xi_n, \\ Y_{n+1} &:= Y_n - h b_n^Y(Y_n) + \sqrt{h} \xi_n, \end{aligned}$$

where $X_0 = Y_0 = x$ and $(\xi_n)_{n \in \mathbb{N}}$ is a sequence of i.i.d. normal $(0, I_d)$ variables. Let $\mathbf{P}, \mathbf{Q}$ denote the joint distributions of $(X_0, X_1, \dots, X_N)$ and $(Y_0, Y_1, \dots, Y_N)$ respectively.

- 1 Write down an expression for the Radon–Nikodym derivative $\frac{d\mathbf{P}}{d\mathbf{Q}}$. Let $h \searrow 0$ and $N \rightarrow \infty$ such that $Nh \rightarrow T$, and compare the result with Girsanov’s theorem (Theorem 3.2.8).
- 2 Write down an expression for $\text{KL}(\mathbf{P} \parallel \mathbf{Q})$. Ensure that it makes sense (i.e., it should be manifestly non-negative).

▷ Exercise 3.4 (Harnack inequalities via shifted Girsanov)

In this exercise, we give another proof of the Harnack inequalities established in Exercise 2.17. Let V be α -convex for some $\alpha \in \mathbb{R}$, and let $(P_t)_{t \geq 0}$ denote the Langevin semigroup with potential V .

- 1 Let $\mathbf{P}'$ be the path measure under which $(B'_t)_{t \geq 0}$ is a standard Brownian motion, and consider the coupled system of SDEs

$$\begin{aligned} dX_t &= \{-\nabla V(X_t) + \eta_t(Y_t - X_t)\} dt + \sqrt{2} dB'_t, & X_0 &= x, \\ dY_t &= -\nabla V(Y_t) dt + \sqrt{2} dB'_t, & Y_0 &= y, \end{aligned}$$

where $(\eta_t)_{t \in [0, T]}$ is a deterministic non-negative process to be chosen later. Assume that we can choose $(\eta_t)_{t \in [0, T]}$ such that $X_T = Y_T$. Finally, let $\mathbf{P}$ be the path measure under which $(B_t)_{t \geq 0}$ is a standard Brownian motion, where

$$dX_t = -\nabla V(X_t) dt + \sqrt{2} dB_t,$$

that is, $dB_t = dB'_t + \frac{\eta_t}{\sqrt{2}}(X_t - Y_t) dt$. Use the data-processing inequality (Theorem 1.5.6) and Girsanov’s theorem (Theorem 3.2.8) to justify the inequality

$$\text{KL}(\delta_y P_T \parallel \delta_x P_T) \leq \text{KL}(\mathbf{P}' \parallel \mathbf{P}) = \frac{1}{4} \mathbb{E}^{\mathbf{P}'} \int_0^T \eta_t^2 \|X_t - Y_t\|^2 dt.$$

2 Use Itô's formula to show that

$$\|X_t - Y_t\|^2 \leq \exp\left(-2\alpha t - 2 \int_0^t \eta_s \, ds\right) \|x - y\|^2.$$

In particular, $X_T = Y_T$ provided $\int_0^T \eta_t \, dt = \infty$. Substitute this into the inequality above and optimize over the choice of $(\eta_t)_{t \in [0, T]}$ in order to establish (2.E.8).

3 Can you extend the argument above in order to establish (2.E.10) as well? *Hint:* Use the fact that the exponential martingale (3.2.7) is a supermartingale.

Doob's Transform

▷ **Exercise 3.5** (Fisher information bound via time reversal)

Let $(X_t)_{t \geq 0}$ denote the Langevin diffusion, let $\mu_t := \text{law}(X_t)$, and let $(X_t^\leftarrow)_{t \in [0, T]}$ denote the time reversal on $[0, T]$.

1 Show that

$$dX_t^\leftarrow = \{\nabla V(X_t^\leftarrow) + 2 \nabla \log \mu_{T-t}(X_t^\leftarrow)\} dt + \sqrt{2} dB_t.$$

2 Using Itô's formula and the Fokker–Planck equation, show that

$$d \nabla \log \frac{\mu_{T-t}}{\pi}(X_t^\leftarrow) = \nabla^2 V(X_t^\leftarrow) \nabla \log \frac{\mu_{T-t}}{\pi}(X_t^\leftarrow) dt + \sqrt{2} \nabla^2 \log \frac{\mu_{T-t}}{\pi}(X_t^\leftarrow) dB_t.$$

Note that while we would naïvely expect Itô's formula to produce a third derivative, there is a miraculous cancellation, due to the time reversal, which removes this term!

3 Using Itô's formula again, show that if $\nabla^2 V \geq \alpha I_d$ for some $\alpha \in \mathbb{R}$, then

$$\text{Fl}(\mu_T \parallel \pi) \leq \exp(-2\alpha T) \text{Fl}(\mu_0 \parallel \pi).$$

Note that this provides a stochastic calculus proof of (2.2.16).

Föllmer Drift

▷ **Exercise 3.6** (Derivation of stochastic localization)

Derive the evolution equation (3.4.3) for stochastic localization.

▷ **Exercise 3.7** (Föllmer and stochastic localization)

Let $\text{BB}_t := (1-t) B_{t/(1-t)}$. By solving the SDE that we derived for Brownian bridge in Example 3.3.4 (with $x_1 = 0$), show that BB is indeed a Brownian bridge. Then, verify that $F_t = (1-t) \theta_{t/(1-t)}$.

▷ **Exercise 3.8** (Kim–Milman map)

In this exercise and the next, we construct a Lipschitz transport map from the standard Gaussian $\gamma := \text{normal}(0, I_d)$ to μ , under appropriate conditions on μ . By Proposition 2.3.3, this implies functional inequalities for μ .

1 For $t \geq 0$, let μ_t denote the law of the Ornstein–Uhlenbeck process at time t , started at $\mu_0 := \mu$. Let $F_t : \mathbb{R}^d \rightarrow \mathbb{R}^d$ denote the flow map at time t , corresponding to the following ODE:

$$\partial_t F_t(x) = -\nabla \log \frac{\mu_t}{\gamma}(F_t(x)), \quad F_0 = \text{id}.$$

Argue that $(F_t)_\# \mu = \mu_t$. *Hint:* Recall the ODE interpretation of the Wasserstein gradient flow, e.g., Exercise 1.13.

- 2 Let $T_t := F_t^{-1}$, so that $(T_t)_\# \mu_t = \mu$. Show that T_t is Lipschitz with

$$\|T_t\|_{\text{Lip}} \leq \exp \int_0^t \lambda_{\max}(\nabla^2 \log \frac{\mu_s}{\gamma}) ds.$$

Thus, we expect that $(T_\infty)_\# \gamma = \mu$, with Lipschitz constant bounded by $\exp \int_0^\infty \lambda_{\max}(\nabla^2 \log(\mu_t/\gamma)) dt$. (This requires justification, but we accept it as a fact here.)

- 3 Establish the identity

$$\nabla^2 \log \frac{\mu_t}{\gamma}(y) = \frac{1}{\exp(2t) - 1} \left(\frac{\text{cov}_{\nu_{1-\exp(-2t), \exp(-t)y}}}{1 - \exp(-2t)} - I_d \right),$$

where for each $\tau > 0$ and $y \in \mathbb{R}^d$ we define the probability measure $\nu_{\tau,y}$ via

$$\nu_{\tau,y}(dx) \propto \exp\left(-\frac{\|x-y\|^2}{2\tau} + \frac{\|x\|^2}{2}\right) \mu(dx).$$

- 4 Suppose that μ is α -strongly log-concave, $\alpha > 1$. (We may assume that the strong log-concavity parameter is sufficiently large, since we can rescale the measure.) Use the identity from the preceding part, together with [Exercise 2.11](#), to compute the integral, and conclude that there exists a Lipschitz transport map from γ to μ with constant at most $\alpha^{-1/2}$.

▷ **Exercise 3.9 (Kim–Milman map II)**

Here, we obtain further conditions for the existence of a Lipschitz transport map from γ to μ .

- 1 Show that if $\text{diam supp } \mu \leq R$, then

$$\nabla^2 \log \frac{\mu_t}{\gamma} \leq \exp(-2t) \left(\frac{R^2}{(1 - \exp(-2t))^2} - \frac{1}{1 - \exp(-2t)} \right) I_d. \quad (3.E.1)$$

- 2 Suppose that $\text{diam supp } \mu \leq R$ and $\nabla^2 \log(1/\mu) \geq \alpha I_d$, where $\alpha \leq 0$. Using a similar estimate as the preceding exercise, write down another bound for $\nabla^2 \log(\mu_t/\gamma)$ which makes use of the assumption $\nabla^2 \log(1/\mu) \geq \alpha I_d$. By taking the minimum of this estimate and (3.E.1), show that there exists a Lipschitz transport map from γ to μ with constant at most $R \exp((1 - \alpha R^2)/2)$. In particular, conclude that if $\alpha = 0$, then μ satisfies a log-Sobolev inequality with constant eR^2 .

- 3 Next, suppose that $\mu = \rho * \gamma$, where $\text{diam supp } \rho \leq R$. Prove that

$$\nabla^2 \log \frac{\mu_t}{\gamma} \leq R^2 \exp(-2t) I_d,$$

and deduce that there is a Lipschitz transport map from γ to μ with constant $\exp(R^2/2)$. Compare this with [Example 2.3.15](#).

- 4 Next, suppose that $\mu \propto \exp(-V - W)$, where V is α -strongly convex ($\alpha > 1$), and W is L -Lipschitz. Argue that

$$\|\text{cov}_{\nu_{\tau,y}}\|_{\text{op}} \leq \left(\sqrt{\|\text{cov}_{\tilde{\nu}_{\tau,y}}\|_{\text{op}}} + W_2(\nu_{\tau,y}, \tilde{\nu}_{\tau,y}) \right)^2,$$

where

$$\tilde{\nu}_{\tau,y}(x) \propto \exp\left(-\frac{\|x-y\|^2}{2\tau} + \frac{\|x\|^2}{2} - V(x)\right),$$

i.e., $\tilde{\nu}_{\tau,y}$ omits the Lipschitz perturbation W . Use [Corollary 3.5.9](#) to control the first term, and the log-Sobolev and T_2 inequalities to control the second term, in order to deduce that

$\|\text{cov}_{\nu_{\tau,y}}\|_{\text{op}} \leq (\sqrt{1/\alpha_\tau} + L/\alpha_\tau)^2$, where $\alpha_\tau := \alpha - 1 + 1/\tau$. Use this to prove that there exists a Lipschitz transport map from γ to μ with constant $\alpha^{-1/2} \exp(2L/\sqrt{\alpha} + L^2/(2\alpha))$. This nearly proves Proposition 2.3.4, albeit with a worse constant; however, we have established a stronger fact, namely the existence of a Lipschitz transport map from a Gaussian.

- 5 Suppose that μ is log-concave, and $\mu_t := \mu * \text{normal}(0, tI_d)$. Prove that there is a 1-Lipschitz map that pushes forward μ_t to μ .

Hint: Use a variant of the Kim–Milman map based on reversing the heat flow, rather than the OU process.

Schrödinger Bridge

▷ Exercise 3.10 (Entropic potentials)

Derive (3.5.4) and use it to prove Lemma 3.5.5 for the entropic potentials.

▷ Exercise 3.11 (Entropic optimal transport between Gaussians)

Compute the entropic optimal transport cost between $\text{normal}(0, 1)$ and $\text{normal}(0, \alpha)$ by making a suitable ansatz for the entropic Brenier potentials.

▷ Exercise 3.12 (Poincaré constant and log-smoothness)

Suppose that $\pi \propto \exp(-V)$, where π satisfies a Poincaré inequality with constant $1/\alpha$, and $\nabla^2 V \preceq \beta I_d$. Show that $\alpha \leq \beta$.

Hint: Use the Poincaré inequality to prove an upper bound on cov_π , and combine it with the lower bound of Corollary 3.5.9.

Part II

Complexity of Sampling

In this chapter, we will provide several analyses of the **Langevin Monte Carlo (LMC)** algorithm, i.e., the iteration

$$\hat{X}_{(n+1)h} := \hat{X}_{nh} - h \nabla V(\hat{X}_{nh}) + \sqrt{2} (B_{(n+1)h} - B_{nh}). \quad (\text{LMC})$$

This is known as the **Euler–Maruyama discretization** of the Langevin diffusion.

Although **LMC** does not achieve state-of-the-art complexity bounds, it is one of the most fundamental sampling algorithms. Through the quantitative convergence analysis of **LMC**, we will develop techniques for discretization analysis that are broadly useful for studying more complex algorithms. To emphasize the kinship of optimization and sampling as the core theme of this book, we also include “optimization boxes” which provide background and context for the corresponding results in optimization theory. Finally, κ always denotes the ratio β/α .

Before proceeding, we state the following fundamental lemma, which will be used repeatedly.

Lemma 4.0.1 (Basic lemma). *Let $\pi \propto \exp(-V)$.*

- 1 *If $\nabla^2 V \geq \alpha I_d > 0$ and V is minimized at $x_\star$, then $\mathbb{E}_\pi[\|\cdot - x_\star\|^2] \leq d/\alpha$.*
- 2 *If $\nabla^2 V \leq \beta I_d$, then $\mathbb{E}_\pi[\|\nabla V\|^2] \leq \beta d$.*
- 3 *If $\nabla^2 V \geq 0$, then $\mathbb{E}_\pi[V - \inf V] \leq d$.*

Proof We use the fact that for the generator $\mathcal{L} = \Delta - \langle \nabla V, \nabla \cdot \rangle$ of the Langevin diffusion, $\mathbb{E}_\pi \mathcal{L} f = 0$ for all test functions $f : \mathbb{R}^d \rightarrow \mathbb{R}$.

- 1 Take $f = \frac{1}{2} \|\cdot - x_\star\|^2$. By strong convexity,

$$0 = \mathbb{E}_\pi \mathcal{L} f = d - \mathbb{E}_\pi \langle \nabla V, \cdot - x_\star \rangle \leq d - \alpha \mathbb{E}_\pi[\|\cdot - x_\star\|^2].$$

- 2 Take $f = V$. Since $\nabla^2 V \leq \beta I_d$, we have $\Delta V \leq \beta d$, whence

$$0 = \mathbb{E}_\pi \mathcal{L} V = \mathbb{E}_\pi[\Delta V - \|\nabla V\|^2] \leq \beta d - \mathbb{E}_\pi[\|\nabla V\|^2].$$

- 3 This follows from differentiating [Lemma 8.6.10](#) at $\lambda = 0$.

□

Stronger sub-Gaussian concentration bounds for ∇V and $V - \inf V$ are given as [Lemma 6.2.7](#) and [Lemma 8.6.10](#) respectively.

4.1 Proof via Wasserstein Coupling

Perhaps the most straightforward analysis of [LMC](#) is based on coupling together the discrete-time algorithm with the continuous-time diffusion, and using this coupling to bound the discretization error in Wasserstein distance. The underlying continuous-time result we use here is the fact that strong log-concavity implies contraction in the Wasserstein metric for the Langevin diffusion. On one hand, this proof is robust and can be applied to more complicated processes; on the other hand, its reliance on contractivity means it is not applicable under weaker assumptions such as an LSI.

Before proceeding, we review the corresponding result for gradient descent.

Optimization Box 4.1.1. Let $V : \mathbb{R}^d \rightarrow \mathbb{R}$ be strongly convex and smooth, i.e., $\alpha I_d \leq \nabla^2 V \leq \beta I_d$. The **gradient descent** (GD) algorithm with fixed step size $h > 0$ is the iteration $x_{n+1} = x_n - h \nabla V(x_n)$. Using strong convexity, we can show that GD converges exponentially fast to the minimizer $x_\star$ of V . First, note that for any $y \in \mathbb{R}^d$, by expanding the square and applying strong convexity,

$$\begin{aligned} \|x_{n+1} - y\|^2 &= \|x_n - h \nabla V(x_n) - y\|^2 \\ &= \|x_n - y\|^2 - 2h \langle \nabla V(x_n), x_n - y \rangle + h^2 \|\nabla V(x_n)\|^2 \\ &\leq (1 - \alpha h) \|x_n - y\|^2 - 2h \{V(x_n) - V(y)\} + h^2 \|\nabla V(x_n)\|^2. \end{aligned}$$

Using the smoothness of V ,

$$V(x_{n+1}) - V(x_n) \leq \langle \nabla V(x_n), x_{n+1} - x_n \rangle + \frac{\beta}{2} \|x_{n+1} - x_n\|^2 = -h \left(1 - \frac{\beta h}{2}\right) \|\nabla V(x_n)\|^2.$$

For any $h \leq 1/\beta$, it yields $\|\nabla V(x_n)\|^2 \leq \frac{2}{h} \{V(x_n) - V(x_{n+1})\}$, and hence

$$\|x_{n+1} - y\|^2 \leq (1 - \alpha h) \|x_n - y\|^2 - 2h \{V(x_{n+1}) - V(y)\}.$$

In particular, for $y = x_\star$, it yields $\|x_{n+1} - x_\star\|^2 \leq (1 - \alpha h) \|x_n - x_\star\|^2$. Choosing $h = 1/\beta$, one obtains $\|x_N - x_\star\| \leq \varepsilon$ in $N = O(\kappa \log(\|x_0 - x_\star\|/\varepsilon))$ iterations, where $\kappa := \beta/\alpha$.

We now consider the corresponding result for [LMC](#).

Theorem 4.1.2 (Coupling analysis of [LMC](#)). *For $n \in \mathbb{N}$, let $\hat{\mu}_{nh}$ denote the law of the n -th iterate of [LMC](#) with step size $h > 0$. Assume that the target $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Then, provided $h \ll 1/\beta\kappa$, for all $N \in \mathbb{N}$,*

$$W_2(\hat{\mu}_{Nh}, \pi) \leq \exp\left(-\frac{\alpha Nh}{2}\right) W_2(\hat{\mu}_0, \pi) + O\left(\frac{\beta d^{1/2} h^{1/2}}{\alpha}\right). \quad (4.1.3)$$

If we set $h \asymp \varepsilon^2/\beta\kappa d$, then for any $\varepsilon \in [0, \sqrt{d}]$ we obtain $\sqrt{\alpha} W_2(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$ after

$$N = O\left(\frac{\kappa^2 d}{\varepsilon^2} \log \frac{\alpha W_2^2(\hat{\mu}_0, \pi)}{\varepsilon^2}\right) \quad \text{iterations}.$$

Remark 4.1.4. We pause to make a few comments about the assumptions and result.

- 1 Typically we assume that we have access to the mode x_* of π , since the complexity of finding the minimizer of V via convex optimization is typically less than the complexity of sampling. In this case, we can initialize at $\mu_0 = \delta_{x_*}$, which satisfies $\alpha W_2^2(\mu_0, \pi) \leq d$ by Lemma 4.0.1.
- 2 It is convenient to use the metric $\sqrt{\alpha} W_2$ instead of W_2 because it is scale-invariant. Namely, for $\lambda > 0$, if we define the scaling map $s_\lambda : x \mapsto \lambda x$, then information divergences such as KL satisfy $\text{KL}((s_\lambda)_\# \mu \parallel (s_\lambda)_\# \pi) = \text{KL}(\mu \parallel \pi)$, by the data-processing inequality (Theorem 1.5.6). On the other hand, W_2 is not invariant, $W_2((s_\lambda)_\# \mu, (s_\lambda)_\# \pi) = \lambda W_2(\mu, \pi)$, but $\sqrt{\alpha} W_2$ is (because the distribution $(s_\lambda)_\# \pi$ is α/λ^2 -strongly log-concave).
Recall also that the T_2 transport inequality, implied by α -strong log-concavity, asserts that $\sqrt{\alpha} W_2(\cdot, \pi) \leq \sqrt{2 \text{KL}(\cdot \parallel \pi)}$. Therefore, $\sqrt{\alpha} W_2$ is a more natural metric.
- 3 The result in (4.1.3) is not sharp; in Section 4.3, via a more sophisticated analysis and averaging, we will improve the iteration complexity to $\tilde{O}(\kappa d/\varepsilon^2)$.
- 4 The inequality (4.1.3) has the following interpretation: for fixed $h > 0$, the first term tends to zero exponentially fast, which reflects the fact that LMC converges to its stationary distribution μ_∞ . However, the stationary distribution is *biased*, $\mu_\infty \neq \pi$, and the second term provides an upper bound on the bias $W_2(\mu_\infty, \pi)$. Note the contrast with Optimization Box 4.1.1, in which there is no bias; this will be discussed further in Section 4.3.

Proof of Theorem 4.1.2 **1. One-step discretization bound.** Suppose that the Langevin diffusion and the LMC algorithm are both initialized at the same measure $\hat{\mu}_0$. We will first bound the discretization error in one step.

We couple the two processes by taking $\hat{X}_0 = X_0$ and using the same Brownian motion:

$$\begin{aligned}\hat{X}_h &= X_0 - h \nabla V(X_0) + \sqrt{2} B_h, \\ X_h &= X_0 - \int_0^h \nabla V(X_t) dt + \sqrt{2} B_h.\end{aligned}$$

Then,

$$\begin{aligned}W_2^2(\hat{\mu}_h, \pi_h) &\leq \mathbb{E}[\|\hat{X}_h - X_h\|^2] \leq \mathbb{E}\left[\left\|\int_0^h \nabla V(X_t) dt - h \nabla V(X_0)\right\|^2\right] \\ &\leq h \int_0^h \mathbb{E}[\|\nabla V(X_t) - \nabla V(X_0)\|^2] dt.\end{aligned}$$

Therefore, we just have to bound the movement $\|X_t - X_0\| = \left\| -\int_0^t \nabla V(X_s) ds + \sqrt{2} B_t \right\|$ of the Langevin diffusion in time t . Roughly, we expect $\left\| \int_0^t \nabla V(X_s) ds \right\| = O(\sqrt{d} t)$ if the size of the gradient is $O(\sqrt{d})$, and $\|B_t\| = O(\sqrt{dt})$. For small t , it is the Brownian motion term which is dominant, which is a common intuition for discretization proofs.

To rigorously bound this term, we appeal to stochastic calculus, see Lemma 4.1.8. For $t \leq 1/3\beta$, it yields the bound

$$\mathbb{E}[\|X_t - X_0\|^2] \leq 8t^2 \mathbb{E}[\|\nabla V(X_0)\|^2] + 8dt$$

and hence, for $\pi_h := \text{law}(X_h)$,

$$W_2^2(\hat{\mu}_h, \pi_h) \leq \beta^2 h \int_0^h \mathbb{E}[\|X_t - X_0\|^2] dt \leq 3\beta^2 h^4 \mathbb{E}[\|\nabla V(X_0)\|^2] + 4\beta^2 dh^3.$$

2. Multi-step discretization bound. We produce a coupling of $\hat{\mu}_{(n+1)h}$ and π as follows. First, let $\hat{X}_{nh} \sim \hat{\mu}_{nh}$ and $X_{nh} \sim \pi$ be optimally coupled. Using the same Brownian motion, we set

$$\begin{aligned}\hat{X}_{(n+1)h} &:= \hat{X}_{nh} - h \nabla V(\hat{X}_{nh}) + \sqrt{2} (B_{(n+1)h} - B_{nh}), \\ X_t &:= X_{nh} - \int_{nh}^t \nabla V(X_s) ds + \sqrt{2} (B_t - B_{nh}), \quad \text{for } t \in [nh, (n+1)h].\end{aligned}$$

Clearly $\hat{X}_{(n+1)h} \sim \hat{\mu}_{(n+1)h}$; also, because π is stationary for the Langevin diffusion, it follows that $X_{(n+1)h} \sim \pi$. We further introduce an auxiliary process: let

$$\bar{X}_t := \hat{X}_{nh} - \int_{nh}^t \nabla V(\bar{X}_s) ds + \sqrt{2} (B_t - B_{nh}) \quad \text{for } t \in [nh, (n+1)h]$$

denote the Langevin diffusion started at $\hat{X}_{nh}$. We bound

$$\begin{aligned}W_2(\hat{\mu}_{(n+1)h}, \pi) &\leq \sqrt{\mathbb{E}[\|\hat{X}_{(n+1)h} - X_{(n+1)h}\|^2]} \\ &\leq \sqrt{\mathbb{E}[\|\bar{X}_{(n+1)h} - X_{(n+1)h}\|^2]} + \sqrt{\mathbb{E}[\|\hat{X}_{(n+1)h} - \bar{X}_{(n+1)h}\|^2]}.\end{aligned}$$

Now we examine the two terms. In the first term, both $\bar{X}$ and X evolve via the Langevin diffusion for an α -strongly convex potential, so we have the following contraction (which is established by a direct coupling argument, see [Theorem 1.4.12](#)):

$$\mathbb{E}[\|\bar{X}_{(n+1)h} - X_{(n+1)h}\|^2] \leq \exp(-2\alpha h) \mathbb{E}[\|\hat{X}_{nh} - X_{nh}\|^2] = \exp(-2\alpha h) W_2^2(\hat{\mu}_{nh}, \pi).$$

For the second term, $\hat{X}$ is the LMC algorithm and $\bar{X}$ is the continuous-time Langevin diffusion, both initialized at the same distribution $\hat{\mu}_{nh}$. Hence, we can apply our one-step discretization bound from before and deduce that

$$\begin{aligned}\mathbb{E}[\|\hat{X}_{(n+1)h} - \bar{X}_{(n+1)h}\|^2] &\leq 3\beta^2 h^4 \mathbb{E}[\|\nabla V(\hat{X}_{nh})\|^2] + 4\beta^2 dh^3 \\ &\leq \beta^4 h^4 \mathbb{E}[\|\hat{X}_{nh} - X_{nh}\|^2] + \beta^2 h^4 \mathbb{E}_\pi[\|\nabla V\|^2] + \beta^2 dh^3 \\ &\leq \beta^4 h^4 W_2^2(\hat{\mu}_{nh}, \pi) + \beta^3 dh^4 + \beta^2 dh^3,\end{aligned}$$

where we used the basic lemma ([Lemma 4.0.1](#)). Note that the second term can be dropped since we are assuming $h \ll 1/\beta$.

Combining everything and using $\sqrt{x+y} \leq \sqrt{x} + \sqrt{y}$ for $x, y \geq 0$,

$$W_2(\hat{\mu}_{(n+1)h}, \pi) \leq \exp(-\alpha h) W_2(\hat{\mu}_{nh}, \pi) + O(\beta^2 h^2 W_2(\hat{\mu}_{nh}, \pi) + \beta d^{1/2} h^{3/2}).$$

Provided $h \ll 1/\beta\kappa$, we can ensure that $\exp(-\alpha h) + O(\beta^2 h^2) \leq \exp(-\alpha h/2)$, therefore absorbing the extra Wasserstein error term. It yields

$$W_2(\hat{\mu}_{(n+1)h}, \pi) \leq \exp\left(-\frac{\alpha h}{2}\right) W_2(\hat{\mu}_{nh}, \pi) + O(\beta d^{1/2} h^{3/2}).$$

After iterating this recursion, it implies

$$W_2(\hat{\mu}_{Nh}, \pi) \leq \exp\left(-\frac{\alpha Nh}{2}\right) W_2(\hat{\mu}_0, \pi) + O\left(\frac{\beta d^{1/2} h^{1/2}}{\alpha}\right).$$

This finishes the proof. $\square$

Remark 4.1.5. By inspecting the proof, one can see that the following stronger inequality holds. Let us denote by $\hat{P}^{\text{LMC}}$ the transition kernel for one step of **LMC**, and P the transition kernel for the Langevin diffusion run for time h . Then, under the assumptions of [Theorem 4.1.2](#), for any $x, y \in \mathbb{R}^d$,

$$W_2(\hat{P}^{\text{LMC}}(x, \cdot), P(y, \cdot)) \leq \exp\left(-\frac{\alpha h}{2}\right) \|x - y\| + O(\beta h^2 \|\nabla V(y)\| + \beta d^{1/2} h^{3/2}). \quad (4.1.6)$$

If we square this inequality and apply Young's inequality, it implies

$$\begin{aligned} W_2^2(\hat{P}^{\text{LMC}}(x, \cdot), P(y, \cdot)) &\leq \exp(-\alpha h) \|x - y\|^2 + O(\beta^2 h^4 \|\nabla V(y)\|^2 + \beta^2 d h^3) \\ &\quad + O(\|x - y\| (\beta h^2 \|\nabla V(y)\| + \beta d^{1/2} h^{3/2})) \\ &\leq \exp\left(-\frac{\alpha h}{2}\right) \|x - y\|^2 + O\left(\frac{\beta^2 h^3 \|\nabla V(y)\|^2}{\alpha} + \frac{\beta^2 d h^2}{\alpha}\right). \end{aligned} \quad (4.1.7)$$

Iterating either (4.1.6) or (4.1.7), together with a coupling argument, recovers the guarantee of [Theorem 4.1.2](#).

We finish by presenting the lemma we used in the proof of [Theorem 4.1.2](#). The following proof is very typical of stochastic calculus arguments, so it is worth internalizing.

Lemma 4.1.8 (Langevin movement bound). *Let $(X_t)_{t \geq 0}$ denote the Langevin diffusion and let $(\pi_t)_{t \geq 0}$ denote its law. Assume that ∇V is β -Lipschitz. Then, provided that $t \leq 1/3\beta$,*

$$\mathbb{E}[\|X_t - X_0\|^2] \leq 8t^2 \mathbb{E}[\|\nabla V(X_0)\|^2] + 8dt.$$

Proof By definition,

$$\mathbb{E}[\|X_t - X_0\|^2] = \mathbb{E}\left[\left\| -\int_0^t \nabla V(X_s) ds + \sqrt{2} B_t \right\|^2\right] \leq 2t \int_0^t \mathbb{E}[\|\nabla V(X_s)\|^2] ds + 4 \mathbb{E}[\|B_t\|^2].$$

Using the β -Lipschitzness of ∇V , $\|\nabla V(X_s)\| \leq \|\nabla V(X_0)\| + \beta \|X_s - X_0\|$. Thus,

$$\mathbb{E}[\|X_t - X_0\|^2] \leq 4\beta^2 t \int_0^t \mathbb{E}[\|X_s - X_0\|^2] ds + 4t^2 \mathbb{E}[\|\nabla V(X_0)\|^2] + 4dt.$$

Applying Grönwall's inequality ([Lemma 1.2.20](#)), it implies

$$\mathbb{E}[\|X_t - X_0\|^2] \leq \{4t^2 \mathbb{E}[\|\nabla V(X_0)\|^2] + 4dt\} \exp(4\beta^2 t^2).$$

Finally, use the assumption $t \leq 1/3\beta$ to conclude. □

4.2 Proof via Interpolation Argument

We now give a guarantee for **LMC** that holds even when V is possibly non-convex.

Optimization Box 4.2.1. Suppose $V : \mathbb{R}^d \rightarrow \mathbb{R}$ is β -smooth but non-convex. Recall from [Optimization Box 4.1.1](#) that along the iterates of GD,

$$V(x_{n+1}) - V(x_n) \leq -\frac{h}{2} \|\nabla V(x_n)\|^2, \quad (4.2.2)$$

provided that $h \leq 1/\beta$ (we only used smoothness to derive this inequality). This is known as the **descent lemma**. We now combine this with an assumption that V satisfies a **Polyak–Łojasiewicz (PL) inequality**

$$\|\nabla V(x)\|^2 \geq 2\alpha \{V(x) - V(x_\star)\} \quad \text{for all } x \in \mathbb{R}^d.$$

The PL inequality is implied by α -strong convexity (see [Section 1.4.2](#)), but it is weaker and allows for non-convex V . We then obtain

$$\begin{aligned} V(x_{n+1}) - V(x_\star) &= V(x_{n+1}) - V(x_n) + V(x_n) - V(x_\star) \\ &\leq -\frac{h}{2} \|\nabla V(x_n)\|^2 + V(x_n) - V(x_\star) \leq (1 - \alpha h) \{V(x_n) - V(x_\star)\}. \end{aligned}$$

Setting $h = 1/\beta$, we can achieve $V(x_N) - V(x_\star) \leq \varepsilon^2$ in $O(\kappa \log(V(x_0) - V(x_\star)/\varepsilon^2))$ iterations under PL and smoothness.

Recall from [Section 1.4.2](#) that the sampling analogue of the PL inequality is the log-Sobolev inequality (LSI), which naturally raises the question of whether the LSI is enough to obtain sampling guarantees. The next proof we give is from [Vempala and Wibisono \(2019\)](#), with a slight refinement using a lemma from [Chewi et al. \(2025a\)](#). Here, we mimic the continuous-time convergence proof in KL divergence by first defining a continuous-time interpolation of the LMC iterates. Upon differentiating the KL divergence along this interpolation, we discover two terms: the first is the Fisher information, and the second is a discretization error term. By controlling the latter, we prove a convergence result for LMC assuming only that π satisfies LSI and that ∇V is Lipschitz.

The interpolation of LMC is defined as follows: we set

$$\hat{X}_t := \hat{X}_{t_-} - (t - t_-) \nabla V(\hat{X}_{t_-}) + \sqrt{2} (B_t - B_{t_-}), \quad (4.2.3)$$

where we have introduced the notation $t_- := \lfloor t/h \rfloor h$.

Proposition 4.2.4. Let $(\hat{\mu}_t)_{t \geq 0}$ be the law of the interpolated process (4.2.3). Then,

$$\partial_t \hat{\mu}_t = \operatorname{div} \left[\hat{\mu}_t \left(\nabla \log \frac{\hat{\mu}_t}{\pi} + \mathbb{E}[\nabla V(\hat{X}_{t_-}) - \nabla V(\hat{X}_t) \mid \hat{X}_t = \cdot] \right) \right].$$

Proof Let φ be a smooth test function. By Itô's formula, (4.2.3) implies

$$\partial_t \mathbb{E} \varphi(\hat{X}_t) = \mathbb{E}[-\langle \nabla \varphi(\hat{X}_t), \nabla V(\hat{X}_{t_-}) \rangle + \Delta \varphi(\hat{X}_t)] = \mathbb{E}[-\langle \nabla \varphi(\hat{X}_t), \mathbb{E}[\nabla V(\hat{X}_{t_-}) \mid \hat{X}_t] \rangle + \Delta \varphi(\hat{X}_t)].$$

Thus,

$$\int \varphi \partial_t \hat{\mu}_t = \partial_t \int \varphi d\hat{\mu}_t = \int \{-\langle \nabla \varphi, \mathbb{E}[\nabla V(\hat{X}_{t_-}) \mid \hat{X}_t = \cdot] \rangle + \Delta \varphi\} d\hat{\mu}_t,$$

and the result follows from integrating by parts. $\square$

Corollary 4.2.5. *Along the law $(\hat{\mu}_t)_{t \geq 0}$ of the interpolated process (4.2.3),*

$$\partial_t \text{KL}(\hat{\mu}_t \parallel \pi) \leq -\frac{3}{4} \text{FI}(\hat{\mu}_t \parallel \pi) + \mathbb{E}[\|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_t) \|^2].$$

Recall that the Fisher information is $\text{FI}(\hat{\mu} \parallel \pi) := \mathbb{E}_{\hat{\mu}}[\|\nabla \log(\hat{\mu}/\pi)\|^2]$ if $\hat{\mu}$ has a smooth density with respect to π .

Proof Using Proposition 4.2.4,

$$\begin{aligned} \partial_t \text{KL}(\hat{\mu}_t \parallel \pi) &= -\mathbb{E}_{\hat{\mu}_t} \left\langle \nabla \log \frac{\hat{\mu}_t}{\pi}, \nabla \log \frac{\hat{\mu}_t}{\pi} + \mathbb{E}[\nabla V(\hat{X}_t) - \nabla V(\hat{X}_t) \mid \hat{X}_t = \cdot] \right\rangle \\ &= -\text{FI}(\hat{\mu}_t \parallel \pi) + \mathbb{E}_{\hat{\mu}_t} \left\langle \nabla \log \frac{\hat{\mu}_t}{\pi}, \mathbb{E}[\nabla V(\hat{X}_t) - \nabla V(\hat{X}_t) \mid \hat{X}_t = \cdot] \right\rangle. \end{aligned}$$

Using Young's inequality,

$$\begin{aligned} &\mathbb{E}_{\hat{\mu}_t} \left\langle \nabla \log \frac{\hat{\mu}_t}{\pi}, \mathbb{E}[\nabla V(\hat{X}_t) - \nabla V(\hat{X}_t) \mid \hat{X}_t = \cdot] \right\rangle \\ &\leq \frac{1}{4} \text{FI}(\hat{\mu}_t \parallel \pi) + \mathbb{E}[\|\mathbb{E}[\nabla V(\hat{X}_t) - \nabla V(\hat{X}_t) \mid \hat{X}_t = \cdot]\|^2] \\ &\leq \frac{1}{4} \text{FI}(\hat{\mu}_t \parallel \pi) + \mathbb{E}[\|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_t) \|^2]. \quad \square \end{aligned}$$

In the convergence proof for LMC, we will use the following lemma. It can be seen as an extension of the second part of the basic lemma (Lemma 4.0.1).

Lemma 4.2.6 (Chewi et al. (2025a, Lemma 16)). *Suppose that $\pi \propto \exp(-V)$ where $\nabla^2 V \leq \beta I_d$. Then, for any probability measure $\mu \ll \pi$,*

$$\mathbb{E}_{\mu}[\|\nabla V\|^2] \leq \text{FI}(\mu \parallel \pi) + 2\beta d.$$

Proof For the generator $\mathcal{L}$ of the Langevin diffusion (with potential V), we can calculate $\mathcal{L}V = \Delta V - \|\nabla V\|^2$. Also, since $\nabla^2 V \leq \beta I_d$, we have $\Delta V \leq \beta d$. Thus, using the fundamental integration by parts identity (Theorem 1.2.14),

$$\begin{aligned} \mathbb{E}_{\mu}[\|\nabla V\|^2] &= \mathbb{E}_{\mu}[\Delta V - \mathcal{L}V] \leq \beta d + \int (-\mathcal{L}V) \frac{d\mu}{d\pi} d\pi = \beta d + \int \langle \nabla V, \nabla \frac{d\mu}{d\pi} \rangle d\pi \\ &= \beta d + \int \langle \nabla V, \nabla \log \frac{d\mu}{d\pi} \rangle d\mu \\ &\leq \beta d + \frac{1}{2} \mathbb{E}_{\mu}[\|\nabla V\|^2] + \frac{1}{2} \text{FI}(\mu \parallel \pi). \end{aligned}$$

Rearranging the inequality yields the result. $\square$

Finally, recall that an LSI implies $\text{KL}(\cdot \parallel \pi) \leq \frac{C_{\text{LSI}}}{2} \text{FI}(\cdot \parallel \pi)$. In the following theorem, we note that $\kappa := \beta/\alpha \geq 1$ as a consequence of Exercise 3.12.

Theorem 4.2.7 (Vempala and Wibisono (2019)). For $n \in \mathbb{N}$, let $\hat{\mu}_{nh}$ denote the law of the n -th iterate of LMC with step size $h > 0$. Assume that the target $\pi \propto \exp(-V)$ satisfies $C_{\text{LSI}}(\pi) \leq 1/\alpha$ and that ∇V is β -Lipschitz. Then, for all $h \leq 1/4\beta$, for all $N \in \mathbb{N}$,

$$\text{KL}(\hat{\mu}_{Nh} \parallel \pi) \leq \exp(-\alpha Nh) \text{KL}(\hat{\mu}_0 \parallel \pi) + O\left(\frac{\beta^2 dh}{\alpha}\right).$$

In particular, if $\kappa := \beta/\alpha$, for all $\varepsilon \in [0, \sqrt{\kappa d}]$ and for $h \asymp \varepsilon^2/\beta\kappa d$, we obtain the guarantee $\sqrt{\text{KL}(\hat{\mu}_{Nh} \parallel \pi)} \leq \varepsilon$ after

$$N = O\left(\frac{\kappa^2 d}{\varepsilon^2} \log \frac{\text{KL}(\hat{\mu}_0 \parallel \pi)}{\varepsilon^2}\right) \quad \text{iterations}.$$

Proof In light of Corollary 4.2.5, we focus our attention on the discretization error term

$$\begin{aligned} \mathbb{E}[\|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_{t-})\|^2] &\leq \beta^2 \mathbb{E}[\|\hat{X}_t - \hat{X}_{t-}\|^2] \\ &\leq \beta^2 h^2 \mathbb{E}[\|\nabla V(\hat{X}_{t-})\|^2] + 2\beta^2 \mathbb{E}[\|B_t - B_{t-}\|^2]. \end{aligned}$$

In order to apply Lemma 4.2.6, it is more convenient to have $\mathbb{E}[\|\nabla V(\hat{X}_t)\|^2]$ instead of $\mathbb{E}[\|\nabla V(\hat{X}_{t-})\|^2]$. So, we use

$$\mathbb{E}[\|\nabla V(\hat{X}_{t-})\|^2] \leq 2 \mathbb{E}[\|\nabla V(\hat{X}_t)\|^2] + 2 \mathbb{E}[\|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_{t-})\|^2].$$

If $h \leq 1/2\beta$, we can combine this inequality with the previous one and rearrange to obtain

$$\begin{aligned} \mathbb{E}[\|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_{t-})\|^2] &\leq 4\beta^2 h^2 \mathbb{E}[\|\nabla V(\hat{X}_t)\|^2] + 4\beta^2 \mathbb{E}[\|B_t - B_{t-}\|^2] \\ &\leq 4\beta^2 h^2 \mathbb{E}[\|\nabla V(\hat{X}_t)\|^2] + 4\beta^2 dh. \end{aligned}$$

For the first term, we apply Lemma 4.2.6, yielding for $h \leq 1/4\beta$

$$4\beta^2 h^2 \mathbb{E}[\|\nabla V(\hat{X}_t)\|^2] \leq 4\beta^2 h^2 \text{Fl}(\hat{\mu}_t \parallel \pi) + 8\beta^3 dh^2 \leq \frac{1}{4} \text{Fl}(\hat{\mu}_t \parallel \pi) + 2\beta^2 dh.$$

Combining with our differential inequality from Corollary 4.2.5 and LSI,

$$\partial_t \text{KL}(\hat{\mu}_t \parallel \pi) \leq -\frac{1}{2} \text{Fl}(\hat{\mu}_t \parallel \pi) + 6\beta^2 dh \leq -\alpha \text{KL}(\hat{\mu}_t \parallel \pi) + 6\beta^2 dh.$$

This implies that

$$\partial_t [\exp(\alpha t) \text{KL}(\hat{\mu}_t \parallel \pi)] \leq 6\beta^2 dh \exp(\alpha t),$$

and upon integration,

$$\text{KL}(\hat{\mu}_{Nh} \parallel \pi) \leq \exp(-\alpha Nh) \text{KL}(\hat{\mu}_0 \parallel \pi) + O\left(\frac{\beta^2 dh}{\alpha}\right). \quad \square$$

Recall from Theorem 2.2.25 that an LSI implies exponential decay in every Rényi divergence, not just the KL divergence. Working with Rényi divergences of order $q > 1$ introduces substantial new difficulties for the discretization analysis, which is why it is remarkable that the above proof can be adapted to the Rényi case with the introduction of some additional tricks; see Chapter 6.

4.3 Proof via Convex Optimization

Next, we turn to an astonishing proof, due to [Durmus et al. \(2019\)](#), which is inspired by convex optimization. This proof yields the state-of-the-art dependence of **LMC** on the condition number κ of the target π , and in Chapter 10, we will also use it to analyze a “mirror” variant of **LMC**.

Let the target be $\pi = \exp(-V)$ (for this proof, we assume that V is normalized so that $\int \exp(-V) = 1$; this just simplifies the notation but changes neither the algorithm nor the analysis). We now view sampling as the composite optimization problem of minimizing the objective

$$\text{KL}(\mu \parallel \pi) := \underbrace{\int V d\mu}_{=: \mathcal{E}(\mu)} + \underbrace{\int \mu \log \mu}_{=: \mathcal{H}(\mu)},$$

where the two terms are the *energy* and the (negative) *entropy*. Accordingly, we break up the iterates of **LMC** into the steps

$$\begin{aligned}\hat{X}_{nh}^+ &:= \hat{X}_{nh} - h \nabla V(\hat{X}_{nh}), \\ \hat{X}_{(n+1)h} &:= \hat{X}_{nh}^+ + \sqrt{2} (B_{(n+1)h} - B_{nh}).\end{aligned}$$

The first step is simply a deterministic gradient descent update on the function V . If we write $\hat{\mu}_{nh}^+$ for the law of $\hat{X}_{nh}^+$, then in the space of measures one can show that $\hat{\mu}_{nh}^+$ is obtained from $\hat{\mu}_{nh}$ by taking a gradient step for the energy functional $\mathcal{E}$ w.r.t. the Wasserstein geometry.¹ On the other hand, the second step applies the heat flow; in the space of measures, this is a Wasserstein gradient flow for the entropy functional $\mathcal{H}$. Since the gradient descent algorithm is sometimes known as the “forward” method in optimization (as opposed to a proximal step which is the “backward” method), this has led to **LMC** being dubbed the “forward–flow” algorithm.

We refer to [Wibisono \(2018\)](#) for more on this perspective. In particular, it suggests why the **LMC** scheme is biased: the “forward–flow” discretization scheme is biased for optimization as well ([Exercise 4.8](#))! Generally speaking, when we split dynamics into its constituent parts and apply a discretization method to each part, this is known as a *splitting scheme*, and it is the cornerstone of numerical integration. Not all splitting schemes are born equal, however—it requires care to design one that is asymptotically unbiased. Recall from [Exercise 1.15](#) that $\mathcal{E}$ is smooth if V is smooth, but $\mathcal{H}$ is non-smooth. Wisdom from optimization theory tells us that the appropriate scheme to use here is the “forward–backward” or “proximal gradient” method (see [Optimization Box 8.0.1](#)), but unfortunately the “backward” step for the entropy cannot be easily implemented.

This can be viewed as the blessing and curse of sampling. It is a blessing because for the non-smooth term in sampling—namely, the entropy $\mathcal{H}$ —although we can implement neither the forward nor the backward discretizations, we *can* implement the exact gradient flow, by sampling a Gaussian, and the gradient flow thankfully succeeds even in the presence of non-smoothness. On the other hand, it is a curse because the “mismatch” of the forward and flow operations as a splitting scheme leads to asymptotic bias. We will revisit this issue of bias in Chapter 8 via the proximal sampler.

Nevertheless, we will leverage this splitting perspective to provide another analysis of **LMC**. The strategy of the proof is to show that the forward step of **LMC** dissipates the energy while not increasing the entropy too much, and that the flow step of **LMC** dissipates the entropy while not increasing the energy too much.

¹Technically this is only true if the step size h is chosen so that $h \|\nabla^2 V\|_{\text{op}} \leq 1$. This is because on any Riemannian manifold in which geodesics cannot be extended indefinitely, the gradient descent steps must be short enough to ensure that the iterates are still travelling along geodesics.

Optimization Box 4.3.1. Let $V : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfy $\alpha I_d \leq \nabla^2 V \leq \beta I_d$, and recall from [Optimization Box 4.1.1](#) that along GD,

$$\|x_{n+1} - y\|^2 \leq (1 - \alpha h) \|x_n - y\|^2 - 2h \{V(x_n) - V(y)\} + h^2 \|\nabla V(x_n)\|^2. \quad (4.3.2)$$

Moreover, for $h \leq 1/\beta$,

$$\|x_{n+1} - y\|^2 \leq (1 - \alpha h) \|x_n - y\|^2 - 2h \{V(x_{n+1}) - V(y)\}. \quad (4.3.3)$$

Inspired by [Ambrosio et al. \(2008\)](#), we refer to this as an **evolution variational inequality (EVI)**. It is quite a flexible tool for optimization. For example, we can set $y = x_\star$, and if $\alpha > 0$, use the fact that $V(x_{n+1}) - V(x_\star) \geq 0$ to conclude that $\|x_{n+1} - x_\star\|^2 \leq (1 - \alpha h) \|x_n - x_\star\|^2$, recovering [Optimization Box 4.1.1](#). However, even when $\alpha = 0$, we can set $y = x_\star$, $h = 1/\beta$, and telescope this inequality to conclude that

$$V(x_N) - V(x_\star) \leq \frac{1}{N} \sum_{n=0}^{N-1} \{V(x_{n+1}) - V(x_\star)\} \leq \frac{\beta \|x_0 - x_\star\|^2}{2N},$$

where the first inequality follows from the descent lemma.

Smooth case.

In analogy, we aim to prove the following key lemma.

Lemma 4.3.4 (EVI with error for [LMC](#)). *Let $\pi = \exp(-V)$ be the target and assume that $0 \leq \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Let $(\hat{\mu}_{nh})_{n \in \mathbb{N}}$ be the marginal laws of [LMC](#) with step size $h \leq 1/\beta$. Then,*

$$2h \text{KL}(\hat{\mu}_{(n+1)h} \parallel \pi) \leq (1 - \alpha h) W_2^2(\hat{\mu}_{nh}, \pi) - W_2^2(\hat{\mu}_{(n+1)h}, \pi) + 2\beta d h^2. \quad (4.3.5)$$

From this, we deduce the following results.

Theorem 4.3.6 ([Durmus et al. \(2019\)](#)). *Suppose that $\pi = \exp(-V)$ is the target distribution and that V satisfies $0 \leq \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Let $(\hat{\mu}_{nh})_{n \in \mathbb{N}}$ denote the marginal laws of [LMC](#).*

1 (Weakly convex case) *Suppose that $\alpha = 0$. For any $\varepsilon \in [0, \sqrt{d}]$, if we take step size $h \asymp \varepsilon^2/\beta d$, then for the mixture distribution $\bar{\mu}_{Nh} := N^{-1} \sum_{n=1}^N \hat{\mu}_{nh}$ it holds that $\sqrt{\text{KL}(\bar{\mu}_{Nh} \parallel \pi)} \leq \varepsilon$ after*

$$N = O\left(\frac{\beta d W_2^2(\hat{\mu}_0, \pi)}{\varepsilon^4}\right) \quad \text{iterations}.$$

2 (Strongly convex case) *Suppose that $\alpha > 0$ and let $\kappa := \beta/\alpha$ denote the condition number. Then, for any $\varepsilon \in [0, \sqrt{d}]$, with step size $h \asymp \varepsilon^2/\beta d$ we obtain $\sqrt{\alpha} W_2(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$ and $\sqrt{\text{KL}(\bar{\mu}_{Nh, 2Nh} \parallel \pi)} \leq \varepsilon$ after*

$$N = O\left(\frac{\kappa d}{\varepsilon^2} \log \frac{\alpha W_2^2(\hat{\mu}_0, \pi)}{\varepsilon^2}\right) \quad \text{iterations},$$

where $\bar{\mu}_{Nh, 2Nh} := N^{-1} \sum_{n=N+1}^{2N} \hat{\mu}_{nh}$.

Proof We use [Lemma 4.3.4](#).

1 By summing the inequality (4.3.5) and using the convexity of the KL divergence,

$$\text{KL}(\bar{\mu}_{Nh} \parallel \pi) \leq \frac{1}{N} \sum_{n=1}^N \text{KL}(\hat{\mu}_{nh} \parallel \pi) \leq \frac{W_2^2(\hat{\mu}_0, \pi)}{2Nh} + \beta dh.$$

The result follows from our choice of h and N .

2 First, we prove the W_2 guarantee. Using the fact that $\text{KL}(\hat{\mu}_{(n+1)h} \parallel \pi) \geq 0$ and iterating the inequality (4.3.5) we obtain

$$\begin{aligned} W_2^2(\hat{\mu}_{Nh}, \pi) &\leq (1 - \alpha h)^N W_2^2(\hat{\mu}_0, \pi) + 2\beta dh^2 \sum_{n=0}^{N-1} (1 - \alpha h)^n \\ &\leq \exp(-\alpha Nh) W_2^2(\hat{\mu}_0, \pi) + O(\kappa dh). \end{aligned}$$

With our choice of h and N , we obtain $\sqrt{\alpha} W_2(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$.

Next, forget about the previous N iterations of LMC and consider $\hat{\mu}_{Nh}$ to be the new initialization to LMC. Applying the weakly convex result yields the KL guarantee $\sqrt{\text{KL}(\bar{\mu}_{Nh, 2Nh} \parallel \pi)} \leq \varepsilon$. $\square$

We next turn toward the proof of Lemma 4.3.4.

Proof of Lemma 4.3.4 We break the proof into three steps.

1. The forward step dissipates the energy. First, let $Z \sim \pi$ be optimally coupled to $\hat{X}_{nh}$. Then, $\mathcal{E}(\hat{\mu}_{nh}^+) - \mathcal{E}(\pi) = \mathbb{E}[V(\hat{X}_{nh}^+) - V(Z)]$. However, since $\hat{X}_{nh}^+$ is obtained from $\hat{X}_{nh}$ via a gradient descent step on V , we can apply the EVI (4.3.3) to argue that

$$\begin{aligned} \mathcal{E}(\hat{\mu}_{nh}^+) - \mathcal{E}(\pi) &\leq \frac{1}{2h} \mathbb{E}[(1 - \alpha h) \|\hat{X}_{nh} - Z\|^2 - \|\hat{X}_{nh}^+ - Z\|^2] \\ &\leq \frac{1}{2h} \{(1 - \alpha h) W_2^2(\hat{\mu}_{nh}, \pi) - W_2^2(\hat{\mu}_{nh}^+, \pi)\}. \end{aligned} \quad (4.3.7)$$

2. The flow step does not substantially increase the energy. Next, using the β -smoothness of V ,

$$\begin{aligned} \mathcal{E}(\hat{\mu}_{(n+1)h}) - \mathcal{E}(\hat{\mu}_{nh}^+) &= \mathbb{E}[V(\hat{X}_{(n+1)h}) - V(\hat{X}_{nh}^+)] \\ &\leq \mathbb{E}[\langle \nabla V(\hat{X}_{nh}^+), \hat{X}_{(n+1)h} - \hat{X}_{nh}^+ \rangle + \frac{\beta}{2} \|\hat{X}_{(n+1)h} - \hat{X}_{nh}^+\|^2] \\ &= \mathbb{E}[\sqrt{2} \langle \nabla V(\hat{X}_{nh}^+), B_{(n+1)h} - B_{nh} \rangle + \beta \|B_{(n+1)h} - B_{nh}\|^2] \\ &= \beta dh. \end{aligned} \quad (4.3.8)$$

3. The flow step dissipates the entropy. Let $(Q_t)_{t \geq 0}$ denote the heat semigroup, i.e., $Q_t f(x) := \mathbb{E} f(x + \sqrt{2} B_t)$, so that $\hat{\mu}_{(n+1)h} = \hat{\mu}_{nh}^+ Q_h$. Then, since the heat flow is the Wasserstein gradient flow of $\mathcal{H}$, and the Wasserstein gradient of $\mathcal{H}$ is $\nabla_{W_2} \mathcal{H}(\mu) = \nabla \log \mu$, one can show that

$$\partial_t W_2^2(\hat{\mu}_{nh}^+ Q_t, \pi) \leq 2 \mathbb{E} \langle \nabla \log \hat{\mu}_{nh}^+ Q_t(\hat{X}_{nh+t}^+), Z - \hat{X}_{nh+t}^+ \rangle$$

where $\hat{X}_{nh+t}^+ \sim \hat{\mu}_{nh}^+ Q_t$ and $Z \sim \pi$ are optimally coupled. This follows from the formula for the gradient of the squared Wasserstein distance (Theorem 1.4.6); it may be justified more rigorously using, e.g., Ambrosio et al. (2008, Theorem 10.2.2).

On the other hand, we showed that $\mathcal{H}$ is geodesically convex in (1.4.4), so

$$\mathcal{H}(\pi) - \mathcal{H}(\hat{\mu}_{nh}^+ Q_t) \geq \mathbb{E} \langle \nabla \log \hat{\mu}_{nh}^+ Q_t(\hat{X}_{nh+t}^+), Z - \hat{X}_{nh+t}^+ \rangle.$$

Using the fact that $t \mapsto \mathcal{H}(\hat{\mu}_{nh}^+ Q_t)$ is decreasing (which also follows because $t \mapsto \hat{\mu}_{nh}^+ Q_t$ is the gradient flow of $\mathcal{H}$), we then have

$$W_2^2(\hat{\mu}_{(n+1)h}, \pi) - W_2^2(\hat{\mu}_{nh}^+, \pi) \leq 2h \{ \mathcal{H}(\pi) - \mathcal{H}(\hat{\mu}_{(n+1)h}) \}. \quad (4.3.9)$$

Concluding the proof. Combine (4.3.7), (4.3.8), and (4.3.9) to obtain (4.3.5). $\square$

Non-smooth case.

The proof via convex optimization can also handle the *non-smooth* case in which we only assume that V is convex and Lipschitz. As before, we deduce the convergence result from a key one-step inequality.

Lemma 4.3.10 (EVI with error for LMC, non-smooth case). *Let $\pi = \exp(-V)$ be the target and assume that V is convex and L -Lipschitz. Let $(\hat{\mu}_{nh})_{n \in \mathbb{N}}$ denote the marginal laws of LMC with step size $h > 0$. Then,*

$$2h \text{KL}(\hat{\mu}_{(n+1)h} \parallel \pi) \leq W_2^2(\hat{\mu}_{nh}^+, \pi) - W_2^2(\hat{\mu}_{(n+1)h}^+, \pi) + L^2 h^2.$$

Theorem 4.3.11 (Durmus et al. (2019)). *Suppose that $\pi = \exp(-V)$ is the target distribution and that V is convex and L -Lipschitz. Let $(\hat{\mu}_{nh})_{n \in \mathbb{N}}$ denote the marginal laws of LMC. For any $\varepsilon > 0$, if we take step size $h \asymp \varepsilon^2 / L^2$, then for the mixture distribution $\bar{\mu}_{Nh} := N^{-1} \sum_{n=1}^N \hat{\mu}_{nh}$ it holds that $\sqrt{\text{KL}(\bar{\mu}_{Nh} \parallel \pi)} \leq \varepsilon$ after*

$$N = O\left(\frac{L^2 W_2^2(\hat{\mu}_0^+, \pi)}{\varepsilon^4}\right) \quad \text{iterations}.$$

Proof This follows from Lemma 4.3.10 in exactly the same way that the weakly convex case of Theorem 4.3.6 follows from Lemma 4.3.4. $\square$

Proof of Lemma 4.3.10 The main task here is to obtain dissipation of the energy functional $\mathcal{E}$ under our new assumptions. Let $Z \sim \pi$ be optimally coupled to $\hat{X}_{(n+1)h}$. From (4.3.2),

$$2h \{ \mathcal{E}(\hat{\mu}_{(n+1)h}) - \mathcal{E}(\pi) \} \leq \mathbb{E}[\|\hat{X}_{(n+1)h} - Z\|^2 - \|\hat{X}_{(n+1)h}^+ - Z\|^2 + h^2 \|\nabla V(\hat{X}_{(n+1)h})\|^2].$$

We no longer have the descent lemma (which requires smoothness) at our disposal, but we can instead bound the last term by $L^2 h^2$ using the Lipschitz assumption. Hence,

$$2h \{ \mathcal{E}(\hat{\mu}_{(n+1)h}) - \mathcal{E}(\pi) \} \leq W_2^2(\hat{\mu}_{(n+1)h}, \pi) - W_2^2(\hat{\mu}_{(n+1)h}^+, \pi) + L^2 h^2. \quad (4.3.12)$$

On the other hand, recall from (4.3.9) that

$$2h \{ \mathcal{H}(\hat{\mu}_{(n+1)h}) - \mathcal{H}(\pi) \} \leq W_2^2(\hat{\mu}_{nh}^+, \pi) - W_2^2(\hat{\mu}_{(n+1)h}, \pi).$$

Together with (4.3.12), this completes the proof. $\square$

4.4 Proof via Girsanov's Theorem

The idea behind the next proof is to control the discretization error $\text{KL}(\hat{\mu}_{Nh} \parallel \pi_{Nh})$, where $\hat{\mu}_{Nh}$ is the law of the LMC iterate $\hat{X}_{Nh}$ and π_{Nh} is the law of the Langevin diffusion X_{Nh} initialized at $\hat{\mu}_0$. This is accomplished using Girsanov's theorem (see Section 3.2.2).

Unlike the previous proofs, which assumed strong log-concavity or LSI for π , controlling this discretization error will only require mild assumptions, namely the smoothness of V . The point, however, is that control of $\text{KL}(\hat{\mu}_{Nh} \parallel \pi_{Nh})$ does not immediately yield quantitative convergence of $\hat{\mu}_{Nh} \rightarrow \pi$; indeed, this also requires quantitative convergence of $\pi_{Nh} \rightarrow \pi$, which does require some kind of assumption such as strong log-concavity or a functional inequality.

Another remark is that the preceding proofs all had the following structure: for a fixed step size $h > 0$, as the number of iterations $N \rightarrow \infty$, the error of **LMC** is at most a quantity depending on h , and which can be made as small as we like by taking h small. In particular, to achieve a desired error it suffices that the step size be sufficiently small and that the number of iterations be sufficiently large. In contrast, for the following proof we will only be able to establish a bound on $\text{KL}(\hat{\mu}_{Nh} \parallel \pi_{Nh})$ which grows with the iteration number N . Consequently, in our final sampling guarantee, we will not be able to take N too large; our guarantee will only imply that the error of **LMC** is small if N lies in some range. Conceptually, this is unsatisfying because “running the Markov chain too long” should not be a problem, and it only arises as an artefact of the proof. Nonetheless, it is worthwhile learning the proof because it is broadly applicable. We remark that this shortcoming can be removed in the strongly log-concave setting using a “shifted Girsanov” approach, similar to [Exercise 3.4](#); see [Altschuler and Chewi \(2026b\)](#) for details.

Historically, the Girsanov method was one of the first discretization techniques utilized in the modern quantitative study of sampling (see [Dalalyan and Tsybakov, 2012](#)). The argument we present here is similar in spirit to [Dalalyan and Tsybakov \(2012\)](#), although we have made a few refinements.

Theorem 4.4.1 (**LMC** discretization bound). *Let $\pi \propto \exp(-V)$ be the target and assume that ∇V is β -Lipschitz. Let $(\hat{\mu}_{nh})_{n \in \mathbb{N}}$ denote the marginal laws of **LMC**, and let $(\pi_t)_{t \geq 0}$ denote the law of the Langevin diffusion initialized at $\hat{\mu}_0$. Then, for $h \lesssim 1/\beta$ and $T := Nh$,*

$$\text{KL}(\hat{\mu}_T \parallel \pi_T) \lesssim \beta^2 h^3 \sum_{n=0}^{N-1} \mathbb{E}_{\hat{\mu}_{nh}} [\|\nabla V\|^2] + \beta^2 dhT .$$

Proof Let $(B_t)_{t \in [0, T]}$ be our standard Brownian motion and consider the SDE

$$dX_t = -\nabla V(X_t) dt + \sqrt{2} dB_t, \quad X_0 \sim \hat{\mu}_0 .$$

Let $\mathbf{W}_T$ denote the Wiener measure on our path space, under which $(X_t)_{t \in [0, T]}$ becomes the Langevin diffusion started at $X_0 \sim \hat{\mu}_0$. We would like to write

$$dX_t = -\nabla V(X_{t-}) dt + \sqrt{2} d\tilde{B}_t ,$$

and to find a path measure $\mathbf{P}_T$ under which $(\tilde{B}_t)_{t \in [0, T]}$ is a $\mathbf{P}_T$ -Brownian motion. If so, then under $\mathbf{P}_T$, we see that $(X_t)_{t \in [0, T]}$ is the interpolated **LMC** process. Noting that $d\tilde{B}_t = dB_t - d[B, M]_t$ where $dM_t := \frac{1}{\sqrt{2}} \langle \nabla V(X_t) - \nabla V(X_{t-}), dB_t \rangle$, consider the exponential martingale $\mathcal{E}(M)$ associated with M .

By the data-processing inequality ([Theorem 1.5.6](#)), $\text{KL}(\hat{\mu}_T \parallel \pi_T) \leq \text{KL}(\mathbf{P}_T \parallel \mathbf{W}_T)$, so it suffices to

bound the latter. By Girsanov's theorem (Theorem 3.2.8),²

$$\begin{aligned} \text{KL}(\mathbf{P}_T \parallel \mathbf{W}_T) &= \mathbb{E}^{\mathbf{P}_T} \log \frac{d\mathbf{P}_T}{d\mathbf{W}_T} = \mathbb{E}^{\mathbf{P}_T} \log \mathcal{E}(M)_T \\ &= \mathbb{E}^{\mathbf{P}_T} \left[\frac{1}{\sqrt{2}} \int_0^T \langle \nabla V(X_t) - \nabla V(X_{t-}), dB_t \rangle - \frac{1}{4} \int_0^T \|\nabla V(X_t) - \nabla V(X_{t-})\|^2 dt \right]. \end{aligned}$$

However, we must be cautious! Here, $(B_t)_{t \in [0, T]}$ is a $\mathbf{W}_T$ -standard Brownian. As a sanity check, if $(B_t)_{t \in [0, T]}$ were a $\mathbf{P}_T$ -standard Brownian motion, then the first term would vanish (since stochastic integrals have zero mean) and the KL divergence would be *negative*, which is absurd. Instead, we rewrite the above expression in terms of $\tilde{B}$, obtaining

$$\begin{aligned} \text{KL}(\mathbf{P}_T \parallel \mathbf{W}_T) &= \frac{1}{4} \int_0^T \mathbb{E}^{\mathbf{P}_T} [\|\nabla V(X_t) - \nabla V(X_{t-})\|^2] dt \leq \frac{\beta^2}{4} \int_0^T \mathbb{E}^{\mathbf{P}_T} [\|X_t - X_{t-}\|^2] dt \\ &= \frac{\beta^2}{4} \int_0^T \{(t - t_-)^2 \mathbb{E}^{\mathbf{P}_T} [\|\nabla V(X_{t-})\|^2] + 2d(t - t_-)\} dt \\ &\leq \frac{\beta^2 h^3}{12} \sum_{n=0}^{N-1} \mathbb{E}^{\mathbf{P}_T} [\|\nabla V(X_{nh})\|^2] + \frac{\beta^2 d h^2 N}{4}. \end{aligned} \quad \square$$

What sampling guarantee does Theorem 4.4.1 imply? Unfortunately, neither KL nor $\sqrt{\text{KL}}$ satisfy the triangle inequality, which poses a difficulty for bounding the distance of $\hat{\mu}_{Nh}$ from the target π . One way to skirt this difficulty is to simply use the fact that $\sqrt{\text{KL}} \gtrsim \|\cdot\|_{\text{TV}}$ (Pinsker's inequality) and the fact that the total variation distance satisfies the triangle inequality.

Corollary 4.4.2. *Let $\pi \propto \exp(-V)$ be the target, assume that $C_{\text{LSI}}(\pi) \leq 1/\alpha$ and that ∇V is β -Lipschitz, and let $\kappa := \beta/\alpha$. Let $(\hat{\mu}_{nh})_{n \in \mathbb{N}}$ denote the marginal laws of LMC. Then, for all sufficiently small $\varepsilon > 0$ and for an appropriate choice of h , we obtain the guarantee $\|\hat{\mu}_{Nh} - \pi\|_{\text{TV}} \leq \varepsilon$ provided that we take*

$$N = \Theta\left(\frac{\kappa^2 d}{\varepsilon^2} \log^2 \frac{\text{KL}(\hat{\mu}_0 \parallel \pi)}{\varepsilon^2}\right) \quad \text{iterations}.$$

For the Pinsker approach, we can leverage the fact that we can bound $\|\hat{\mu}_{Nh} - \pi_{Nh}\|_{\text{TV}}^2$ either by $\text{KL}(\hat{\mu}_{Nh} \parallel \pi_{Nh})$ or $\text{KL}(\pi_{Nh} \parallel \hat{\mu}_{Nh})$, depending on which is easier. The latter KL divergence is controlled in exactly the same way as in Theorem 4.4.1, except that in the resulting bound we replace $\mathbb{E}_{\hat{\mu}_{nh}}[\|\nabla V\|^2]$ with $\mathbb{E}_{\pi_{nh}}[\|\nabla V\|^2]$.

Proof of Corollary 4.4.2 The LSI implies exponential convergence of the Langevin diffusion to its target in KL divergence (Theorem 1.2.26 and Theorem 1.2.30), and we obtain $\sqrt{\text{KL}(\pi_{Nh} \parallel \pi)} \leq \varepsilon/\sqrt{2}$ with the choice $T = Nh \asymp \alpha^{-1} \log(\text{KL}(\hat{\mu}_0 \parallel \pi)/\varepsilon^2)$.

Next, if we simply apply the discretization bound for $\text{KL}(\hat{\mu}_{Nh} \parallel \pi_{Nh})$ in Theorem 4.4.1, it becomes tricky to bound $\mathbb{E}_{\hat{\mu}_{nh}}[\|\nabla V\|^2]$ without further assumptions. Instead, as discussed above, we use a bound for $\text{KL}(\pi_{Nh} \parallel \hat{\mu}_{Nh})$. We apply Talagrand's T_2 inequality, which is implied by the LSI (Exercise 1.16),

²Actually, as noted in the discussion in Section 3.2.2, to obtain the first equality one should check Novikov's condition. However, since all we desire is an upper bound on the KL divergence, this can be avoided with a localization argument. We omit the details.

and the data-processing inequality³ (Theorem 1.5.6):

$$\mathbb{E}_{\pi_{nh}}[\|\nabla V\|^2] \lesssim \mathbb{E}_{\pi}[\|\nabla V\|^2] + \beta^2 W_2^2(\pi_{nh}, \pi) \lesssim \beta d + \frac{\beta^2}{\alpha} \text{KL}(\pi_{nh} \parallel \pi) \leq \beta d + \frac{\beta^2}{\alpha} \text{KL}(\hat{\mu}_0 \parallel \pi).$$

Similarly to Theorem 4.4.1, we then obtain

$$\text{KL}(\pi_{Nh} \parallel \hat{\mu}_{Nh}) \lesssim \beta^3 dh^2 T + \frac{\beta^4 h^2 T}{\alpha} \text{KL}(\hat{\mu}_0 \parallel \pi) + \beta^2 dh T.$$

We now choose h to make this at most $\varepsilon^2/2$. To simplify the calculation, we assume that ε is sufficiently small so that the third term is dominant, as is typically the case in applications. This leads to the choice $h = \Theta(\varepsilon^2 / \beta \kappa d \log(\text{KL}(\hat{\mu}_0 \parallel \pi) / \varepsilon^2))$.

By the triangle inequality and Pinsker's inequality,

$$\|\hat{\mu}_{Nh} - \pi\|_{\text{TV}} \leq \|\hat{\mu}_{Nh} - \pi_{Nh}\|_{\text{TV}} + \|\pi_{Nh} - \pi\|_{\text{TV}} \leq \sqrt{\frac{1}{2} \text{KL}(\pi_{Nh} \parallel \hat{\mu}_{Nh})} + \sqrt{\frac{1}{2} \text{KL}(\pi_{Nh} \parallel \pi)} \leq \varepsilon.$$

Finally, plugging in the choice of h into $Nh \asymp \alpha^{-1} \log(\text{KL}(\hat{\mu}_0 \parallel \pi) / \varepsilon^2)$ yields the result. $\square$

Although the quantitative dependence in Corollary 4.4.2 matches prior results (e.g., via the interpolation method in Theorem 4.2.7), the final result is unsatisfying because we have moved to a weaker metric (TV rather than KL) for a seemingly silly reason (the failure of the triangle inequality for the KL divergence). Indeed, we have a convergence result for the Langevin diffusion in KL, and our discretization bound is in KL, yet our final result is in TV. Can we remedy this?

To address this, we can introduce the Rényi divergences, defined in (2.2.24); recall that $\text{KL} = \mathcal{R}_1$. We have a continuous-time result for the Langevin diffusion in Rényi divergence (Theorem 2.2.25), and it turns out that with some additional tricks it is possible to extend the Girsanov discretization argument to any Rényi divergence. Moreover, the Rényi divergences satisfy a *weak triangle inequality* (see Lemma 6.2.5). This allows us to combine a continuous-time Rényi result with a Rényi discretization argument to yield a Rényi sampling guarantee. We provide the details for this approach in Chapter 6.

Bibliographical Notes

Historically, the LMC algorithm, which is called *unadjusted* because of the lack of a Metropolis–Hastings filter, was only studied relatively recently in non-asymptotic settings. Before the work of Dalalyan and Tsybakov (2012), it was more common to study MALA (which we introduce and study in Chapter 7), partly due to the influence of Roberts and Tweedie (1996). The ideas that go into the basic W_2 coupling proof for Theorem 4.1.2 were developed in a series of works on strongly log-concave sampling: Dalalyan and Tsybakov (2012); Dalalyan (2017a,b); Durmus and Moulines (2017, 2019). The Girsanov argument of Theorem 4.4.1 is also due to Dalalyan and Tsybakov (2012).

The coupling analysis of Section 4.1 will be generalized into a local error framework in Section 5.1.1. Using this framework, Exercise 5.2 will show that the rate for LMC can be improved under higher-order smoothness assumptions on V . Guarantees in KL divergence under higher-order smoothness were also obtained via the interpolation method of Section 4.2 in Mou et al. (2022).

The proof of Exercise 4.7 is taken from Altschuler and Chewi (2024b), and the examples of Exercise 4.8 and Exercise 4.9 are from Wibisono (2018).

³The data-processing inequality is not available to us if we bound the KL divergence with the opposite order of arguments.

A coupling-based approach to the proof in Section 4.3 was given in Dalalyan and Karagulyan (2026). They also observed that a refined notion of smoothness can be used in the analysis (Exercise 4.10).

Another notable proof technique that we have omitted from this chapter is **reflection coupling** (Eberle, 2011, 2016). Reflection coupling uses a carefully chosen coupling of the Brownian motions rather than just taking the two Brownian motions to be the same as we have done (the latter coupling is called the **synchronous coupling**).

Exercises

Proof via Wasserstein Coupling

⊢ Exercise 4.1 (Explicit computations for a Gaussian target)

Suppose that the target distribution is a Gaussian, $\pi = \text{normal}(0, \Sigma)$, and that LMC is initialized at a Gaussian. Can you write down the iterates and stationary distribution of LMC explicitly? What happens when $\Sigma = I_d$?

Perform some explicit computations for this example and compare them to the general results for LMC that we derived in this chapter. In particular, for $\Sigma = I_d$, show that the step size can be taken to be $h \asymp \varepsilon/d^{1/2}$ to control the bias.

⊢ Exercise 4.2 (Second moment bounds for LMC)

Assume that $\beta I_d \geq \nabla^2 V \geq \alpha I_d > 0$ and $\nabla V(0) = 0$. Prove that if $h \leq 1/\beta$, then the LMC iterates initialized at δ_0 with step size $h > 0$ have uniformly bounded second moment: $\sup_{n \in \mathbb{N}} \mathbb{E}[\|\hat{X}_{nh}\|^2] \leq d/\alpha$. (Write a recursion for $\mathbb{E}[\|\hat{X}_{(n+1)h}\|^2]$ in terms of $\mathbb{E}[\|\hat{X}_{nh}\|^2]$. In order to prove the result for all step sizes $h \leq 1/\beta$, you may need to appeal to coercivity of the gradient, Lemma 5.2.3.)

⊢ Exercise 4.3 (LMC with decaying step size)

Consider LMC initialized at $\delta_{x_\star}$, where $x_\star$ is the mode of π and π satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Show that by considering LMC with a decaying step size schedule $h_n = 1/C\beta\kappa + can$, where $C, c > 0$ are appropriately chosen universal constants, one can obtain an iteration complexity which removes the logarithmic factor in Theorem 4.1.2 for all $N \gtrsim \kappa^2$.

Hint: The term $C\beta\kappa$ in the choice of step size is so that the one-step bound in Remark 4.1.5 is valid. However, the resulting computation is tedious, so here is a simpler exercise that still conveys the main point. Take the step size schedule $h_n = 1/can$ for $c > 0$ sufficiently small and use the one-step bound from Remark 4.1.5 (pretending that it is still valid) to obtain an iteration complexity bound without a logarithmic term.

⊢ Exercise 4.4 (Uniform-in-time LSI for LMC under strong log-concavity)

Here, we give an alternative, discrete-time approach to Theorem 2.2.21.

- 1 Let $(\hat{\mu}_{nh})_{n \in \mathbb{N}}$ denote the laws of the iterates of LMC, and assume that $\alpha I_d \leq \nabla^2 V \leq \beta I_d$, where α is allowed to be negative. Show that if h is small enough (depending on β only), and if $C_{\text{LSI}}(\mu_0) \leq 1/\alpha_0$, then for all $n \in \mathbb{N}$,

$$C_{\text{LSI}}(\hat{\mu}_{nh}) \leq \frac{c_1(\alpha, h, n)}{\alpha_0} + c_2(\alpha, h, n),$$

for some $c_1(\alpha, h, n)$ and $c_2(\alpha, h, n)$ that you should compute. In particular, if $\alpha > 0$, then the LSI is uniform-in-time: $\sup_{n \in \mathbb{N}} C_{\text{LSI}}(\hat{\mu}_{nh}) < \infty$.

Hint: Use Proposition 2.3.3 and Proposition 2.3.8.

- 2 Take the continuous-time limit $h \searrow 0$ as $nh \rightarrow t$ in the previous part, and deduce that the local LSI (Theorem 2.2.21) holds.

Proof via Interpolation Argument

⊢ **Exercise 4.5** (More general interpolations)

Let $(X_t)_{t \geq 0}$ be a stochastic process with C^1 sample paths, and let $\mu_t := \text{law}(X_t)$. Prove that

$$\partial_t \mu_t + \text{div}(\mu_t b_t) = 0,$$

where $b_t := \mathbb{E}[\dot{X}_t \mid X_t = \cdot]$. Note that this Fokker–Planck equation also describes the marginal laws of the process $\dot{Y}_t = b_t(Y_t)$, which is Markovian. Hence, it is known as a *Markovian projection*; see Gyöngy (1986); Brunick and Shreve (2013).

More generally, for any $\gamma > 0$, show that $\partial_t \mu_t + \text{div}(\mu_t b_t^\gamma) = \gamma \Delta \mu_t$, where $b_t^\gamma := b_t + \gamma \nabla \log \mu_t$. This corresponds to the SDE $dY_t = b_t^\gamma(Y_t) dt + \sqrt{2\gamma} dB_t$.

⊢ **Exercise 4.6** (Fokker–Planck vs. Girsanov's theorem)

Consider the Fokker–Planck equations for two SDEs:

$$\partial_t \mu_t + \text{div}(\mu_t b_t) = \gamma \Delta \mu_t, \quad \partial_t \bar{\mu}_t + \text{div}(\bar{\mu}_t \bar{b}_t) = \gamma \Delta \bar{\mu}_t.$$

By differentiating $t \mapsto \text{KL}(\mu_t \parallel \bar{\mu}_t)$ (similarly to the calculation in Corollary 4.2.5) show that $\text{KL}(\mu_T \parallel \bar{\mu}_T) \leq \frac{1}{4\gamma} \int_0^T \mathbb{E}_{\mu_t} [\|b_t - \bar{b}_t\|^2] dt$. Compare this with what can be obtained from Girsanov's theorem (Theorem 3.2.8).

⊢ **Exercise 4.7** (Harnack inequalities via Fokker–Planck)

In this exercise, we provide yet another proof of the Harnack inequalities from Exercise 2.17 and Exercise 3.4. Consider Fokker–Planck equations for the Langevin diffusion and for an auxiliary process, given as follows:

$$\partial_t \mu_t = \text{div}(\mu_t \nabla \log \frac{\mu_t}{\pi}), \quad \partial_t \nu_t = \text{div}(\nu_t \nabla \log \frac{\nu_t}{\pi}), \quad \partial_t \mu'_t = \text{div}(\mu'_t (\nabla \log \frac{\mu'_t}{\pi} - \eta_t (T_{\mu'_t \rightarrow \mu_t} - \text{id})))$$

with $\mu'_0 = \nu_0$ and $(\eta_t)_{t \geq 0}$ chosen so that $\mu'_T = \mu_T$. Here, $T_{\mu'_t \rightarrow \mu_t}$ is the optimal transport map from μ'_t to μ_t . Assuming that $\nabla^2 V \geq \alpha I_d$, differentiate $t \mapsto \text{KL}(\mu'_t \parallel \nu_t)$ as in the preceding exercise, and use this to again establish (2.E.8). Can you also adapt this argument to Rényi divergences?

Proof via Convex Optimization

⊢ **Exercise 4.8** (Forward–flow discretization is biased)

Consider optimization of a composite objective $f+g$, where $f(x) := \frac{1}{2}(x-1)^2$ and $g(x) := \frac{1}{2}(x+1)^2$. Write down the “forward–flow” splitting scheme for minimizing $f+g$, compute the limiting point of the algorithm with step size h , and deduce that this scheme is biased.

⊢ **Exercise 4.9** (Proximal operator for the entropy)

As described further in Optimization Box 8.0.1, the sampling analogue of a “backward” discretization step on the entropy would be to solve the following optimization problem:

$$\text{prox}_{h\mathcal{H}}(\nu) := \arg \min_{\mu \in \mathcal{P}_2(\mathbb{R}^d)} \left\{ \mathcal{H}(\mu) + \frac{1}{2h} W_2^2(\mu, \nu) \right\}.$$

In general, this is intractable. However, show that when $\nu = \text{normal}(0, \Sigma)$, the solution μ is also Gaussian, and determine the mean and covariance of μ .

▷ **Exercise 4.10** (Refined smoothness)

Show that the bound in (4.3.8) can be replaced with $\beta_\Delta dh$, where $\beta_\Delta := d^{-1} \sup \Delta V \leq \beta$. What final iteration complexity does this imply for LMC?

▷ **Exercise 4.11** (LMC without a two-phase analysis)

In Theorem 4.3.6, we deduced the rate for the strongly convex case via a two-phase analysis; in particular, we throw away the first N iterates and average the following N iterates, which can be interpreted as a “burn-in” phase. Here, we show that this is unnecessary, provided that we use a different averaging scheme. Starting with (4.3.5), multiply both sides by $(1 - \alpha h)^{-(n+1)}$ and sum across the iterations. Show that a suitable average $\bar{\mu}_{Nh}$ of the first iterates achieves the KL guarantee in Theorem 4.3.6 for the strongly convex case.

Proof via Girsanov’s Theorem

▷ **Exercise 4.12** (Convergence under PI)

Follow the proof of Corollary 4.4.2, but instead of assuming that $C_{\text{LSI}}(\pi) \leq 1/\alpha$, assume instead that $C_{\text{PI}}(\pi) \leq 1/\alpha$. What is the final iteration complexity if $\log \chi^2(\hat{\mu}_0 \parallel \pi) = \tilde{O}(d)$? (See Lemma 1.5.2.)

▷ **Exercise 4.13** (Diverging Girsanov bounds)

Consider the case where $\pi = \text{normal}(0, I_d)$ and show that for the path measures in the proof of Theorem 4.4.1, $\text{KL}(\mathbf{P}_T \parallel \mathbf{W}_T)$ necessarily grows linearly with T ; thus, the diverging bound is a fundamental limitation of the Girsanov-based approach. See, however, Altschuler and Chewi (2026b) for how to circumvent this with the shifted Girsanov method. On the other hand, argue that $\text{KL}(\hat{\mu}_T \parallel \pi_T)$ remains bounded as $T \rightarrow \infty$, so the data-processing inequality is a bit loose here.

▷ **Exercise 4.14** (Sharper Girsanov bound)

Show that in Theorem 4.4.1, the following sharper upper bound holds:

$$\text{KL}(\hat{\mu}_T \parallel \pi_T) \leq \frac{1}{4} \int_0^T \mathbb{E} \left[\left\| \mathbb{E}[\nabla V(\hat{X}_{t-}) \mid \hat{X}_t] - \nabla V(\hat{X}_t) \right\|^2 \right] dt.$$

Thus, Girsanov’s theorem can also capture some of the benefits of the interpolation method. *Hint:* Similarly to Exercise 4.5, use Proposition 4.2.4 to argue that there is a Markov process which is different from the interpolation of LMC, yet has the same marginal laws as LMC. Apply Girsanov’s theorem to this process instead.

We now move beyond the basic [LMC](#) algorithm and consider samplers with better dependence on the dimension and inverse accuracy. There are two main sources of improvement that we explore in this chapter. The first is to use a more sophisticated discretization method than the basic Euler–Maruyama discretization. The second is to consider a different stochastic process, called the *underdamped* Langevin diffusion. By combining these two ideas, we arrive at state-of-the-art complexity bounds for low-accuracy samplers.

5.1 Randomized Midpoint Discretization

In this section, we study the **randomized midpoint discretization**, which was introduced in [Shen and Lee \(2019\)](#). The application to the Langevin diffusion was carried out in [He et al. \(2020\)](#) and more recently in [Yu et al. \(2024\)](#).

Consider the continuous-time Langevin diffusion from time nh to $(n+1)h$:

$$X_{(n+1)h} = X_{nh} - \int_{nh}^{(n+1)h} \nabla V(X_t) dt + \sqrt{2} (B_{(n+1)h} - B_{nh}).$$

In the Euler discretization, we approximate the second term via $-h \nabla V(X_{nh})$. However, if we want an *unbiased* estimator of the integral, then we can introduce an auxiliary random variable $u_n \sim \text{uniform}([0, 1])$ and use

$$X_{(n+1)h} \approx X_{nh} - h \nabla V(X_{nh+u_n h}) + \sqrt{2} (B_{(n+1)h} - B_{nh}).$$

To compute this approximation, however, we need to know $X_{nh+u_n h}$. We have

$$X_{nh+u_n h} = X_{nh} - \int_{nh}^{nh+u_n h} \nabla V(X_t) dt + \sqrt{2} (B_{nh+u_n h} - B_{nh}).$$

Note that we are in the same situation as before. In particular, if we desire, we can draw another uniform random variable u'_n and approximate the second term above via $-u'_n h \nabla V(X_{nh+u_n u'_n h})$. In principle, this procedure can be repeated indefinitely. However, we will see that just one step of this

procedure suffices: further applications of this procedure do not improve the discretization error. Instead, we will simply approximate X_{nh+u_nh} via an Euler–Maruyama step:

$$X_{nh+u_nh} \approx X_{nh} - u_n h \nabla V(X_{nh}) + \sqrt{2} (B_{nh+u_nh} - B_{nh}).$$

To summarize, the randomized midpoint discretization of the Langevin diffusion, which we will call RM–LMC, is the following update:

$$\begin{aligned} \hat{X}_{(n+1)h} &:= \hat{X}_{nh} - h \nabla V(\hat{X}_{nh+u_nh}^+) + \sqrt{2} (B_{(n+1)h} - B_{nh}), \\ \hat{X}_{nh+u_nh}^+ &:= \hat{X}_{nh} - u_n h \nabla V(\hat{X}_{nh}) + \sqrt{2} (B_{nh+u_nh} - B_{nh}), \end{aligned} \tag{RM–LMC}$$

where $(u_n)_{n \in \mathbb{N}}$ is a sequence of i.i.d. uniform $([0, 1])$ random variables which are independent of $\hat{X}_0$ and the Brownian motion. The algorithm uses two gradient evaluations per iteration. Also, when implementing this recursion, it is important to note that the two Brownian increments are *coupled*. To sample the Brownian increments, draw two i.i.d. standard Gaussians ξ_n and ξ'_n , and set

$$\begin{aligned} B_{nh+u_nh} - B_{nh} &:= \sqrt{u_n h} \xi_n, \\ B_{(n+1)h} - B_{nh} &:= \sqrt{u_n h} \xi_n + \sqrt{(1 - u_n)h} \xi'_n. \end{aligned}$$

We now state the complexity bound for RM–LMC.

Theorem 5.1.1 (RM–LMC). *For $n \in \mathbb{N}$, let $\hat{\mu}_{nh}$ denote the law of the n -th iterate of RM–LMC with step size $h > 0$. Assume that the target $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Then, provided $h \ll 1/\beta\kappa^{1/2}$, for all $N \in \mathbb{N}$,*

$$W_2^2(\hat{\mu}_{Nh}, \pi) \leq \exp\left(-\frac{\alpha Nh}{2}\right) W_2^2(\hat{\mu}_0, \pi) + O\left(\frac{\beta^2 dh^2}{\alpha} + \frac{\beta^4 dh^3}{\alpha^2}\right).$$

In particular, if we choose h appropriately, then for any $\varepsilon \in [0, d^{1/2}/\kappa^{1/4}]$ we obtain the guarantee $\sqrt{\alpha} W_2(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$ after

$$N = \tilde{O}\left(\frac{\kappa d^{1/2}}{\varepsilon} \vee \frac{\kappa^{4/3} d^{1/3}}{\varepsilon^{2/3}}\right) \quad \text{iterations}.$$

This result is a substantial improvement over the $\tilde{O}(\kappa d/\varepsilon^2)$ rate obtained for LMC in Theorem 4.3.6. To establish the theorem, we now introduce a powerful framework known as *local error analysis*, which streamlines the analysis of RM–LMC as well as other sophisticated discretizations of SDEs.

See the Bibliographical Notes for more recent developments.

5.1.1 Local Error Analysis

The intuition behind randomized midpoint is that the gradient evaluated at a uniformly chosen random time is an “unbiased” estimator of the stochastic integral. But what exactly does “bias” mean in this context? Suppose that the “error” incurred by a discretization method at iteration n is denoted ω_n , and for the sake of argument, suppose that $\omega_1, \dots, \omega_N$ are i.i.d. with mean m and variance σ^2 . Then, the total error over N iterations, $\sum_{n=1}^N \omega_n$, has size roughly $Nm + N^{1/2}\sigma$, where the first term represents systematic bias that adds up over the iterations and the second term represents stochastic fluctuations which cancel out and hence only add up as the square root of the number of iterations. Although this model for the errors is clearly unrealistic, it motivates the development of a framework that separates

out the effects of “systematic bias” versus “stochastic fluctuations”, which in the sequel will be called the “weak error” and the “strong error”.

The following framework, known as **local error analysis** or **mean-squared error analysis**, is classical in the literature on numerical methods for SDEs (Milstein and Tretyakov, 2021).

Lemma 5.1.2 (Local error analysis). *Let $\hat{P}, P$ be two Markov kernels over $\mathbb{R}^d$ and let $x, y \in \mathbb{R}^d$. Suppose that there are jointly defined random variables $\hat{X} \sim \delta_x \hat{P}$, $X \sim \delta_x P$, $Y \sim \delta_y P$ satisfying the following four conditions:*

- (Contractivity) $\|X - Y\|_{L^2} \leq L \|x - y\|$.
- (Coupling) $\|(X - x) - (Y - y)\|_{L^2} \leq \zeta \|x - y\|$.
- (Weak error) $\|\mathbb{E} \hat{X} - \mathbb{E} X\| \leq \mathcal{E}_{\text{weak}}(x)$.
- (Strong error) $\|\hat{X} - X\|_{L^2} \leq \mathcal{E}_{\text{strong}}(x)$.

Then, the following one-step Wasserstein bound holds:

$$W_2^2(\delta_x \hat{P}, \delta_y P) \leq L^2 \|x - y\|^2 + 2(\mathcal{E}_{\text{weak}}(x) + \zeta \mathcal{E}_{\text{strong}}(x)) \|x - y\| + \mathcal{E}_{\text{strong}}(x)^2. \quad (5.1.3)$$

In particular, if the above conditions hold for all $x, y \in \mathbb{R}^d$ for some $L < 1$, then for any probability measures μ_0, ν_0 ,

$$W_2^2(\mu_0 \hat{P}^N, \nu_0 P^N) \leq L^N W_2^2(\mu_0, \nu_0) + \frac{(\bar{\mathcal{E}}_{\text{weak}} + \zeta \bar{\mathcal{E}}_{\text{strong}})^2}{L(1-L)^2} + \frac{(\bar{\mathcal{E}}_{\text{strong}})^2}{1-L},$$

where we set

$$\bar{\mathcal{E}}_{\text{weak}} := \max_{n=0,1,\dots,N-1} \|\mathcal{E}_{\text{weak}}\|_{L^2(\mu_0 \hat{P}^n)}, \quad \bar{\mathcal{E}}_{\text{strong}} := \max_{n=0,1,\dots,N-1} \|\mathcal{E}_{\text{strong}}\|_{L^2(\mu_0 \hat{P}^n)}.$$

Proof Expand out the square and apply Cauchy–Schwarz:

$$\begin{aligned} \mathbb{E}[\|\hat{X} - Y\|^2] &= \mathbb{E}[\|X - Y\|^2 + 2\langle \hat{X} - X, X - Y \rangle + \|\hat{X} - X\|^2] \\ &\leq L^2 \|x - y\|^2 + \mathbb{E}[2\langle \hat{X} - X, x - y \rangle + 2\langle \hat{X} - X, (X - x) - (Y - y) \rangle + \|\hat{X} - X\|^2] \\ &\leq L^2 \|x - y\|^2 + 2(\mathcal{E}_{\text{weak}}(x) + \zeta \mathcal{E}_{\text{strong}}(x)) \|x - y\| + \mathcal{E}_{\text{strong}}(x)^2. \end{aligned}$$

To iterate this bound, we apply Young’s inequality to the middle term to obtain

$$2(\mathcal{E}_{\text{weak}}(x) + \zeta \mathcal{E}_{\text{strong}}(x)) \|x - y\| \leq L(1-L) \|x - y\|^2 + \frac{1}{L(1-L)} (\mathcal{E}_{\text{weak}}(x) + \zeta \mathcal{E}_{\text{strong}}(x))^2.$$

Hence,

$$W_2^2(\delta_x \hat{P}, \delta_y P) \leq L \|x - y\|^2 + \frac{1}{L(1-L)} (\mathcal{E}_{\text{weak}}(x) + \zeta \mathcal{E}_{\text{strong}}(x))^2 + \mathcal{E}_{\text{strong}}(x)^2.$$

A coupling argument (cf. Lemma 8.2.2) and iteration yield the result. $\square$

In applications, P is taken to be an idealized process (e.g., the exact Langevin diffusion), whereas $\hat{P}$ is a numerical discretization. The lemma adopts four assumptions, two of which are solely properties of the kernel P . The remaining two assumptions, which quantify the accuracy of the numerical scheme, are purely local in that they only consider one step of the algorithm and idealized process started from the same starting point, and are therefore typically easy to compute.

Note that the effective time horizon arising from a contraction factor L is $N_{\text{eff}} := 1/(1-L)$. So, if

we disregard the term involving ζ (which is often small), the final bound shows that the error over N steps is bounded by $N_{\text{eff}} \bar{\mathcal{E}}_{\text{weak}} + N_{\text{eff}}^{1/2} \bar{\mathcal{E}}_{\text{strong}}$, in accordance with the intuition presented above.

As stated, the local error framework requires contractivity ($L < 1$) and only produces guarantees in the W_2 metric. Both of these points were addressed in [Altschuler and Chewi \(2026b\)](#), which established a local error framework in KL divergence. One advantage of such a framework is that the errors do not grow exponentially with the number of iterations, even in the case $L > 1$.

Remark 5.1.4. It is often more convenient to have the weak and strong error terms depend on y rather than x , because the expectations are then taken under the idealized process $\{\nu_0 P^n\}_{n \geq 0}$. To handle this, often $\mathcal{E}_{\text{weak}}, \mathcal{E}_{\text{strong}}$ are $L_{\text{weak}}, L_{\text{strong}}$ -Lipschitz respectively, in which case (5.1.3) implies

$$W_2^2(\delta_x \hat{P}, \delta_y P) \leq L'^2 \|x - y\|^2 + 2(\mathcal{E}_{\text{weak}}(y) + \zeta \mathcal{E}_{\text{strong}}(y)) \|x - y\| + 2\mathcal{E}_{\text{strong}}(y)^2,$$

with $L'^2 = L^2 + 2L_{\text{weak}} + 2\zeta L_{\text{strong}} + 2L_{\text{strong}}^2$. If $L' < 1$, this can be iterated to yield

$$W_2^2(\mu_0 \hat{P}^N, \nu_0 P^N) \leq L'^N W_2^2(\mu_0, \nu_0) + \frac{(\bar{\mathcal{E}}_{\text{weak}} + \zeta \bar{\mathcal{E}}_{\text{strong}})^2}{L' (1 - L')^2} + \frac{2(\bar{\mathcal{E}}_{\text{strong}})^2}{1 - L'},$$

where now we take

$$\bar{\mathcal{E}}_{\text{weak}} = \max_{n=0,1,\dots,N-1} \|\mathcal{E}_{\text{weak}}\|_{L^2(\nu_0 P^n)}, \quad \bar{\mathcal{E}}_{\text{strong}} = \max_{n=0,1,\dots,N-1} \|\mathcal{E}_{\text{strong}}\|_{L^2(\nu_0 P^n)}.$$

Remark 5.1.5. If we do not wish to take advantage of the weak error, then we instead have the simpler recursion $W_2(\delta_x \hat{P}, \delta_y P) \leq L \|x - y\| + \mathcal{E}_{\text{strong}}(x)$. This is iterated to yield

$$W_2(\mu_0 \hat{P}^N, \nu_0 P^N) \leq L^N W_2(\mu_0, \nu_0) + \frac{\max_{n=0,1,\dots,N-1} \|\mathcal{E}_{\text{strong}}\|_{L^2(\mu_0 \hat{P}^n)}}{1 - L}.$$

If $\mathcal{E}_{\text{strong}}$ is L_{strong} -Lipschitz, then for $L' := L + L_{\text{strong}}$,

$$W_2(\mu_0 \hat{P}^N, \nu_0 P^N) \leq L'^N W_2(\mu_0, \nu_0) + \frac{\max_{n=0,1,\dots,N-1} \|\mathcal{E}_{\text{strong}}\|_{L^2(\nu_0 P^n)}}{1 - L'}.$$

To conclude this section, we show that the Langevin diffusion satisfies the first two assumptions of [Lemma 5.1.2](#).

Lemma 5.1.6 (Local error assumptions for the Langevin diffusion). *Let $(X_t)_{t \geq 0}, (Y_t)_{t \geq 0}$ denote two copies of the Langevin diffusion with potential V and started at x and y respectively. Assume that $0 \leq \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Then, the following statements hold.*

- 1 $\|X_h - Y_h\|_{L^\infty} \leq \exp(-\alpha h) \|x - y\|.$
- 2 $\|X_h - x - (Y_h - y)\|_{L^\infty} \leq \beta h \|x - y\|.$

Proof The first statement is [Exercise 1.17](#). For the second statement,

$$\|X_h - x - (Y_h - y)\| = \left\| \int_0^h \{\nabla V(X_t) - \nabla V(Y_t)\} dt \right\| \leq \beta \int_0^h \|X_t - Y_t\| dt \leq \beta h \|x - y\|. \quad \square$$

5.1.2 Analysis of Randomized Midpoint

To prove [Theorem 5.1.1](#), we compute the weak and strong errors for RM-LMC.

Lemma 5.1.7 (Local errors for RM-LMC). *Let $\hat{P}$, P denote the kernels for RM-LMC and for the Langevin diffusion run for time h respectively. Assume that $-\beta I_d \leq \nabla^2 V \leq \beta I_d$ and $h \ll 1/\beta$. There is a coupling $\hat{X}_h \sim \delta_x \hat{P}$, $X_h \sim \delta_x P$ such that the following hold.*

- 1 (Weak error) $\|\mathbb{E} \hat{X}_h - \mathbb{E} X_h\| \lesssim \beta^2 h^3 \|\nabla V(x)\| + \beta^2 d^{1/2} h^{5/2}$.
- 2 (Strong error) $\|\hat{X}_h - X_h\|_{L^2} \lesssim \beta h^2 \|\nabla V(x)\| + \beta d^{1/2} h^{3/2}$.

Proof We begin with the strong error.

$$\|\hat{X}_h - X_h\| = \left\| \int_0^h \{\nabla V(\hat{X}_{uh}^+) - \nabla V(X_t)\} dt \right\| \leq \beta \int_0^h \|\hat{X}_{uh}^+ - X_t\| dt.$$

Next,

$$\begin{aligned} \|\hat{X}_{uh}^+ - X_t\|_{L^2} &= \left\| \int_0^t \nabla V(X_s) ds - uh \nabla V(x) + \sqrt{2} (B_{uh} - B_t) \right\|_{L^2} \\ &\lesssim h \|\nabla V(x)\| + \int_0^t \|\nabla V(X_s)\|_{L^2} ds + d^{1/2} h^{1/2} \\ &\lesssim h \|\nabla V(x)\| + \beta \int_0^t \|X_s - x\|_{L^2} ds + d^{1/2} h^{1/2} \lesssim h \|\nabla V(x)\| + d^{1/2} h^{1/2}, \end{aligned}$$

where we applied Lemma 4.1.8. This yields the strong error estimate.

For the weak error, note that by taking the expectation over the uniform random variable, we obtain

$$\mathbb{E}[\hat{X}_h \mid (B_t)_{t \in [0, h]}] = x - \int_0^h \nabla V(\hat{X}_t^+) dt + \sqrt{2} B_h.$$

This leads to substantial cancellations:

$$\begin{aligned} \|\mathbb{E} \hat{X}_h - \mathbb{E} X_h\| &= \left\| \mathbb{E} \int_0^h \{\nabla V(\hat{X}_t^+) - \nabla V(X_t)\} dt \right\| \leq \beta \int_0^h \mathbb{E} \|\hat{X}_t^+ - X_t\| dt \\ &= \beta \int_0^h \mathbb{E} \left\| \int_0^t \{\nabla V(X_s) - \nabla V(x)\} ds \right\| dt \leq \beta^2 \int_0^h \int_0^t \mathbb{E} \|X_s - x\| ds dt \\ &\lesssim \beta^2 h^2 \{h \|\nabla V(x)\| + d^{1/2} h^{1/2}\}, \end{aligned}$$

where we again applied Lemma 4.1.8. □

Proof of Theorem 5.1.1 We apply Lemma 5.1.2 via Remark 5.1.4 and Lemma 5.1.6. In particular, for $h \ll 1/\beta\kappa^{1/2}$, we can take $L' = \exp(-\alpha h/2)$. By Lemma 4.0.1, we have the estimates $\tilde{\mathcal{E}}_{\text{weak}} \lesssim \beta^2 d^{1/2} h^{5/2}$ and $\tilde{\mathcal{E}}_{\text{strong}} \lesssim \beta d^{1/2} h^{3/2}$. The result follows after some algebraic simplifications.¹ □

5.2 Hamiltonian Monte Carlo

The next algorithm we introduce, known as **Hamiltonian Monte Carlo (HMC)**, was introduced as “hybrid Monte Carlo” in Duane et al. (1987) and later popularized for statistical sampling by Neal (2011). As the name suggests, it is inspired by Hamiltonian mechanics. Although this algorithm is usually combined with a Metropolis–Hastings filter, we defer a discussion of this until Chapter 7. In this section, we instead focus on an analysis of the ideal (i.e., continuous-time) dynamics.

¹ Check this with paper and pen.

5.2.1 Introduction to Ideal HMC

First, we augment the target distribution π to add a momentum variable p . Specifically, define the distribution π on *phase space* $\mathbb{R}^d \times \mathbb{R}^d$ via

$$\pi(x, p) \propto \exp(-V(x) - \frac{1}{2} \|p\|^2).$$

The first marginal of π is $\pi \propto \exp(-V)$, so a sample from π yields, upon projecting to the first coordinate, a sample from π .

The augmented target can also be written as $\pi \propto \exp(-H)$, where H is the **Hamiltonian** $H(x, p) := V(x) + \frac{1}{2} \|p\|^2$. In Hamiltonian mechanics, which is a reformulation of classical mechanics, the laws of motion are governed by **Hamilton's equations**, a system of coupled first-order ODEs:²

$$\begin{aligned}\dot{x}_t &= \nabla_p H(x_t, p_t) = p_t, \\ \dot{p}_t &= -\nabla_x H(x_t, p_t) = -\nabla V(x_t).\end{aligned}$$

Introducing the anti-symmetric matrix

$$J := \begin{bmatrix} 0 & I_d \\ -I_d & 0 \end{bmatrix},$$

Hamilton's equations can be written succinctly as

$$(\dot{x}_t, \dot{p}_t) = J \nabla H(x_t, p_t).$$

Let $F_t : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}^d \times \mathbb{R}^d$ denote the **flow map**, i.e., $F_t(x_0, p_0)$ is the solution (x_t, p_t) to Hamilton's equations started from (x_0, p_0) . Then, we show that F_t leaves the augmented target π invariant: $(F_t)_\# \pi = \pi$. Indeed, if $f : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$ is a function on phase space and $(x_t, p_t)_{t \geq 0}$ evolve via Hamilton's equations started at $(x_0, p_0) \sim \pi$,

$$\begin{aligned}\partial_t \Big|_{t=0} \mathbb{E} f(x_t, p_t) &= \mathbb{E} \langle \nabla f(x_0, p_0), (\dot{x}_0, \dot{p}_0) \rangle = \int \langle \nabla f(x_0, p_0), (\dot{x}_0, \dot{p}_0) \rangle \pi(dx_0, dp_0) \\ &= - \int f \operatorname{div} \left(\begin{bmatrix} p \\ -\nabla V \end{bmatrix} \pi \right) \\ &= - \int f \left\{ \operatorname{div} \begin{bmatrix} p \\ -\nabla V \end{bmatrix} + \left\langle \begin{bmatrix} p \\ -\nabla V \end{bmatrix}, \begin{bmatrix} \nabla_x \log \pi \\ \nabla_p \log \pi \end{bmatrix} \right\rangle \right\} d\pi = 0.\end{aligned}$$

Further properties of the Hamiltonian dynamics are explored in [Exercise 5.5](#).

However, simply running Hamilton's equations does not yield a convergent sampling algorithm. For example, suppose that $V(x) = \frac{1}{2} \|x\|^2$; then, each flow map F_t is actually a diffeomorphism. This implies, for example, that $\operatorname{KL}((F_t)_\# \mu \parallel (F_t)_\# \pi) = \operatorname{KL}(\mu \parallel \pi)$ for any initial distribution μ on phase space. To get around this issue, we can “refresh” the momentum periodically. More specifically, we pick an integration time $T > 0$, and every T units of time we draw a new momentum vector from the standard Gaussian distribution (which is the distribution of the momentum under π).

Ideal HMC: Pick an integration time $T > 0$ and draw $(X_0, P_0) \sim \mu_0$. For each iteration $n = 0, 1, 2, \dots$:

- 1 *Refresh* the velocity by drawing $P'_{nT} \sim \operatorname{normal}(0, I_d)$.
- 2 *Integrate* Hamilton's equations: set $(X_{(n+1)T}, P_{(n+1)T}) := F_T(X_{nT}, P'_{nT})$.

²In contrast, Newton's law $\ddot{x}_t = -\nabla V(x_t)$ is a second-order ODE.

Since both steps of each iteration preserve π , the entire algorithm preserves π . At this stage, though, this algorithm is still idealized because it assumes the ability to exactly integrate Hamilton's equations. This may be possible for very special cases, but it is not in general, and certainly not within our oracle model. Nevertheless, it is instructive to first analyze the ideal algorithm.

5.2.2 Analysis of Ideal HMC

Theorem 5.2.1 (Ideal HMC, [Chen and Vempala \(2019\)](#)). Assume that the target $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. For $n \in \mathbb{N}$, let μ_{nT} denote the law of the n -th iterate X_{nT} of ideal HMC with integration time $T > 0$. Then, if we set $T = 1/2\sqrt{\beta}$, we obtain

$$W_2^2(\mu_{NT}, \pi) \leq \exp\left(-\frac{N}{16\kappa}\right) W_2^2(\mu_0, \pi).$$

The rate of contraction in this theorem is optimal, see [Chen and Vempala \(2019\)](#). We now follow the proof, which is a purely deterministic analysis of Hamilton's equations. First, we need two lemmas.

Lemma 5.2.2 (A priori bound). Let $(x_t, p_t)_{t \geq 0}$ and $(x'_t, p'_t)_{t \geq 0}$ denote two solutions to Hamilton's equations of motion with $p_0 = p'_0$ and a potential V satisfying $\|\nabla^2 V\|_{\text{op}} \leq \beta$. Then, for all $t \in [0, 1/2\sqrt{\beta}]$, it holds that

$$\frac{1}{2} \|x_0 - x'_0\|^2 \leq \|x_t - x'_t\|^2 \leq 2 \|x_0 - x'_0\|^2.$$

Proof First, note that $\partial_t \|p_t - p'_t\| \leq \|\nabla V(x_t) - \nabla V(x'_t)\| \leq \beta \|x_t - x'_t\|$. It follows that $|\partial_t \|x_t - x'_t\|| \leq \|p_t - p'_t\| \leq \beta \int_0^t \|x_s - x'_s\| ds$ and hence

$$\|x_t - x'_t\| \leq \|x_0 - x'_0\| + \beta \int_0^t \int_0^s \|x_r - x'_r\| dr ds.$$

Applying an ODE comparison lemma, one may deduce that $\|x_t - x'_t\| \leq \sqrt{2} \|x_0 - x'_0\|$. The lower bound is similar. $\square$

Lemma 5.2.3 (Coercivity). Suppose $f : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfies $0 \leq \nabla^2 f \leq \beta I_d$. Then, for all $x, y \in \mathbb{R}^d$, it holds that

$$\beta \langle \nabla f(x) - \nabla f(y), x - y \rangle \geq \|\nabla f(x) - \nabla f(y)\|^2.$$

Proof See [Exercise 5.6](#). $\square$

We now prove the contraction result for the Hamiltonian dynamics.

Proposition 5.2.4 (Contraction of Hamilton's equations). Consider any two solutions $(x_t, p_t)_{t \geq 0}$ and $(x'_t, p'_t)_{t \geq 0}$ to Hamilton's equations of motion with $p_0 = p'_0$ and a potential V satisfying $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Then, for all $t \in [0, 1/2\sqrt{\beta}]$, it holds that

$$\|x_t - x'_t\|^2 \leq \exp\left(-\frac{\alpha t^2}{4}\right) \|x_0 - x'_0\|^2.$$

Proof We compute

$$\begin{aligned} \frac{1}{2} \partial_t \|x_t - x'_t\|^2 &= \langle x_t - x'_t, p_t - p'_t \rangle, \\ \frac{1}{2} \partial_t^2 \|x_t - x'_t\|^2 &= \|p_t - p'_t\|^2 - \langle x_t - x'_t, \nabla V(x_t) - \nabla V(x'_t) \rangle = -\rho_t \|x_t - x'_t\|^2 + \|p_t - p'_t\|^2, \end{aligned}$$

where we define

$$\rho_t := \frac{\langle \nabla V(x_t) - \nabla V(x'_t), x_t - x'_t \rangle}{\|x_t - x'_t\|^2}.$$

To bound $\|p_t - p'_t\|^2$, we use $|\partial_t \|p_t - p'_t\|| \leq \|\nabla V(x_t) - \nabla V(x'_t)\|$. Also, by the coercivity lemma (Lemma 5.2.3),

$$\|\nabla V(x_t) - \nabla V(x'_t)\|^2 \leq \beta \langle \nabla V(x_t) - \nabla V(x'_t), x_t - x'_t \rangle = \beta \rho_t \|x_t - x'_t\|^2 \leq 2\beta \rho_t \|x_0 - x'_0\|^2$$

where we used Lemma 5.2.2. Hence, by the Cauchy–Schwarz inequality,

$$\|p_t - p'_t\|^2 \leq \left| \int_0^t |\partial_s \|p_s - p'_s\|| \, ds \right|^2 \leq \left| \int_0^t \sqrt{2\beta \rho_s} \|x_0 - x'_0\| \, ds \right|^2 \leq 2\beta t \|x_0 - x'_0\|^2 \int_0^t \rho_s \, ds.$$

From this and Lemma 5.2.2, we deduce

$$\partial_t^2 \|x_t - x'_t\|^2 \leq -\left(\rho_t - 4\beta t \int_0^t \rho_s \, ds\right) \|x_0 - x'_0\|^2.$$

Integrating and using $\rho \geq 0$,

$$\begin{aligned} \partial_t \|x_t - x'_t\|^2 &\leq -\left(\int_0^t \rho_s \, ds - 4\beta \int_0^t s \int_0^s \rho_r \, dr \, ds\right) \|x_0 - x'_0\|^2 \\ &\leq -\left(\int_0^t \rho_s \, ds - 2\beta t^2 \int_0^t \rho_s \, ds\right) \|x_0 - x'_0\|^2 \\ &= -(1 - 2\beta t^2) \left(\int_0^t \rho_s \, ds\right) \|x_0 - x'_0\|^2 \leq -\frac{\alpha t}{2} \|x_0 - x'_0\|^2. \end{aligned}$$

Integrating again then yields

$$\|x_t - x'_t\|^2 \leq \|x_0 - x'_0\|^2 - \frac{\alpha t^2}{4} \|x_0 - x'_0\|^2 \leq \exp\left(-\frac{\alpha t^2}{4}\right) \|x_0 - x'_0\|^2. \quad \square$$

If we choose $T = 1/2\sqrt{\beta}$, then the contraction factor is $\exp(-1/16\kappa) \leq 1 - 1/32\kappa$. In particular, let P denote the transition kernel of one step of ideal HMC with integration time T on the position coordinate. We have the W_1 contraction

$$W_1(P(x, \cdot), P(x', \cdot)) \leq \exp\left(-\frac{1}{32\kappa}\right) \|x - x'\|,$$

which, in the language of Section 2.6, says that the coarse Ricci curvature of P is bounded below by $\Omega(\kappa^{-1})$. Hence, by Theorem 2.6.6, we immediately obtain the following corollary.

Corollary 5.2.5. *Assume that the target $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Let P be the Markov kernel for ideal HMC with integration time $T = 1/2\sqrt{\beta}$. Then, P satisfies a Poincaré inequality with constant $O(\kappa)$, where $\kappa := \beta/\alpha$.*

The contractivity result of [Theorem 5.2.1](#) only holds for “short” integration times, up to time of order $1/\sqrt{\beta}$. This is optimal, in the sense that the contraction in [Proposition 5.2.4](#) fails beyond this time horizon. However, it is still possible to establish *accelerated* convergence rates for HMC in a more nuanced, average-case sense, using randomized integration times which can be long—up to order $1/\sqrt{\alpha}$. We summarize the results here.

Theorem 5.2.6 (Accelerated convergence of ideal HMC). *Assume that the target $\pi \propto \exp(-V)$ is log-concave, and let μ_T denote the law of randomized HMC with total (expected) integration time T . There exist choices for the law of the integration times and a universal constant $c > 0$ such that the following results hold for all $T > 0$.*

1 ([Lu and Wang \(2022b\)](#)) Suppose that π satisfies a Poincaré inequality with constant α^{-1} . Then,

$$\chi^2(\mu_T \parallel \pi) \lesssim \exp(-c\sqrt{\alpha}T) \chi^2(\mu_0 \parallel \pi) .$$

2 ([Mitra et al. \(2026\)](#); [Monmarché and Wang \(2026\)](#)) Suppose that π satisfies a log-Sobolev inequality with constant α^{-1} . Then,

$$\text{KL}(\mu_T \parallel \pi) \lesssim \exp(-c\sqrt{\alpha}T) \text{KL}(\mu_0 \parallel \pi) .$$

3 ([Mitra et al. \(2026\)](#)) Under log-concavity alone,

$$\text{KL}(\mu_T \parallel \pi) \lesssim \frac{\text{KL}(\mu_0 \parallel \pi) + W_2^2(\mu_0, \pi)}{T^2} .$$

These accelerated results are quite subtle and we omit their proofs, opting to postpone the more detailed discussion of the acceleration phenomenon to the next section in the context of the underdamped Langevin diffusion. We remark that, at present, the results of [Theorem 5.2.6](#) have not yet been translated into discrete-time implementations.

5.3 The Underdamped Langevin Diffusion

In ideal HMC, the momentum is refreshed periodically. We now consider a variant in which the momentum is refreshed continuously. The **underdamped Langevin diffusion** is the solution to

$$\begin{aligned} dX_t &= P_t dt , \\ dP_t &= -\nabla V(X_t) dt - \gamma P_t dt + \sqrt{2\gamma} dB_t . \end{aligned} \tag{5.3.1}$$

Here, $\gamma > 0$ is a parameter known as the **friction parameter**; as the name suggests, the physical interpretation is that the Hamiltonian system is damped by friction.

The underdamped Langevin diffusion is motivated by the acceleration phenomenon in optimization, which we first recall in continuous time.

Optimization Box 5.3.2. Consider the following ODE system with $p_0 = 0$:

$$\begin{aligned}\dot{x}_t &= p_t, \\ \dot{p}_t &= -\nabla V(x_t) - \gamma p_t.\end{aligned}$$

This is the deterministic analogue of the underdamped Langevin diffusion. If we assume that V is α -strongly convex, then one can show that with the choice $\gamma = 2\sqrt{\alpha}$,

$$V(x_t) - V(x_\star) \leq 2 \exp(-\sqrt{\alpha} t) \{V(x_0) - V(x_\star)\}, \quad (5.3.3)$$

see [Exercise 5.7](#). Moreover, the ODE system is stable for integration times of order $t \asymp 1/\sqrt{\beta}$, where β is the smoothness of V , and hence one expects discretization of this system to yield an algorithm for optimization with a *square root* dependence on the condition number κ . This is indeed the case, but discretization is subtle; see, e.g., [Nesterov \(2018, Section 2.2\)](#) for a discrete-time scheme which achieves $V(x_N) - V(x_\star) \leq \varepsilon^2$ in $O(\sqrt{\kappa} \log(\kappa (V(x_0) - V(x_\star))/\varepsilon^2))$ iterations.

This phenomenon is known as **acceleration** in optimization. Historically, the development happened in the opposite order: Nesterov put forth his algorithm, now known as **Nesterov's accelerated gradient descent**, in [Nesterov \(1983\)](#), and his algorithm is optimal amongst all first-order algorithms ([Nemirovsky and Yudin, 1983](#)). The continuous-time formulation of acceleration was introduced later, in [Su et al. \(2016\)](#).

Motivated by the acceleration phenomenon in optimization, we undertake a detailed study of the underdamped Langevin diffusion to see if such a phenomenon also holds for log-concave sampling. At present, however, our understanding is incomplete.

Unlike the Langevin diffusion, the underdamped Langevin diffusion is *not* a reversible Markov process. Moreover, it is an example of non-trivial *hypocoercive dynamics*, which means that the Markov semigroup approach based on Poincaré and log-Sobolev inequalities fails, necessitating the use of more sophisticated PDE analysis (see [Section 5.3.4](#)).

5.3.1 Continuous-Time Considerations

It is illuminating to write the dynamics in the space of measures. By computing the generator of this Markov process and writing down the corresponding Fokker–Planck equation, we arrive at the PDE

$$\partial_t \pi_t = \gamma \Delta_p \pi_t + \operatorname{div} \left(\pi_t \begin{bmatrix} -p \\ \nabla V + \gamma p \end{bmatrix} \right) \quad (5.3.4)$$

for the evolution of the law $(\pi_t)_{t \geq 0}$ of the underdamped Langevin diffusion. We can write this as the continuity equation

$$\partial_t \pi_t = \operatorname{div} \left(\pi_t \begin{bmatrix} \nabla_p \log \pi \\ -\nabla_x \log \pi - \gamma \nabla_p \log \pi + \gamma \nabla_p \log \pi_t \end{bmatrix} \right)$$

where $\pi(x, p) \propto \exp(-H(x, p)) = \exp(-V(x) - \frac{1}{2} \|p\|^2)$. However, taking advantage of the fact that

$$\operatorname{div} \left(\pi_t \begin{bmatrix} -\nabla_p \log \pi_t \\ \nabla_x \log \pi_t \end{bmatrix} \right) = 0,$$

we have the more interpretable expression

$$\partial_t \pi_t = \operatorname{div} \left(\pi_t J_\gamma \begin{bmatrix} \nabla_x \log(\pi_t / \pi) \\ \nabla_p \log(\pi_t / \pi) \end{bmatrix} \right), \quad J_\gamma := \begin{bmatrix} 0 & -1 \\ 1 & \gamma \end{bmatrix} \otimes I_d,$$

or

$$\partial_t \pi_t = \operatorname{div} \left(\pi_t J_\gamma [\nabla_{W_2} \operatorname{KL}(\cdot \| \pi)](\pi_t) \right). \quad (5.3.5)$$

This shows that the underdamped Langevin diffusion is not interpreted as a gradient flow of the KL divergence, but rather a “damped Hamiltonian flow” for the KL divergence.

We begin with a contraction result for the continuous-time process based on [Cheng et al. \(2018\)](#). Note that we use the same change of variables as in [Exercise 5.7](#).

Theorem 5.3.6 (Contraction of underdamped Langevin). *Let $(X_t^0, P_t^0)_{t \geq 0}$ and $(X_t^1, P_t^1)_{t \geq 0}$ be two copies of the underdamped Langevin diffusion, driven by the same Brownian motion. Assume that the potential satisfies $0 \leq \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Then, defining the modified norm*

$$\| (x, p) \|_{\gamma}^2 := \|x\|^2 + \left\| x + \frac{2}{\gamma} p \right\|^2$$

and for all $\gamma \geq \sqrt{2\beta}$, we obtain the contraction

$$\| (X_t^1, P_t^1) - (X_t^0, P_t^0) \|_{\gamma} \leq \exp\left(-\frac{\alpha t}{\gamma}\right) \| (X_0^1, P_0^1) - (X_0^0, P_0^0) \|_{\gamma}.$$

Proof Write $\delta X_t := X_t^1 - X_t^0$ and $\delta P_t := P_t^1 - P_t^0$. Then, by Itô’s formula ([Theorem 1.1.19](#)),

$$\begin{aligned} d\left(\delta X_t + \frac{2}{\gamma} \delta P_t\right) &= \left[\delta P_t - \frac{2}{\gamma} \{\nabla V(X_t^1) - \nabla V(X_t^0)\} - 2\delta P_t\right] dt \\ &= \left[-\delta P_t - \frac{2}{\gamma} \underbrace{\left(\int_0^1 \nabla^2 V((1-s)X_t^0 + sX_t^1) ds\right)}_{=: H_t} \delta X_t\right] dt \\ &= \left[-\frac{\gamma}{2} \left(\delta X_t + \frac{2}{\gamma} \delta P_t\right) + \left(\frac{\gamma}{2} - \frac{2}{\gamma} H_t\right) \delta X_t\right] dt \end{aligned}$$

as well as

$$d(\delta X_t) = \delta P_t dt = \left[\frac{\gamma}{2} \left(\delta X_t + \frac{2}{\gamma} \delta P_t\right) - \frac{\gamma}{2} \delta X_t\right] dt$$

so that

$$\begin{aligned} &\frac{1}{2} \partial_t \left\{ \left\| \delta X_t + \frac{2}{\gamma} \delta P_t \right\|^2 + \|\delta X_t\|^2 \right\} \\ &= - \left\langle \begin{bmatrix} \delta X_t + \frac{2}{\gamma} \delta P_t \\ \delta X_t \end{bmatrix}, \begin{bmatrix} \gamma/2 I_d & 1/2 (2/\gamma H_t - \gamma I_d) \\ 1/2 (2/\gamma H_t - \gamma I_d) & \gamma/2 I_d \end{bmatrix} \begin{bmatrix} \delta X_t + (2/\gamma) \delta P_t \\ \delta X_t \end{bmatrix} \right\rangle. \end{aligned}$$

One checks that for $\gamma \geq \sqrt{2\beta}$, the eigenvalues of the matrix above are lower bounded by α/γ . $\square$

We check that the new norm we defined is equivalent to the Euclidean norm.

Lemma 5.3.7 (Equivalence of norms). *For all $x, p \in \mathbb{R}^d$ and $\gamma > 0$,*

$$\frac{1}{3} \left(\|x\|^2 + \frac{4}{\gamma^2} \|p\|^2 \right) \leq \| (x, p) \|^2 \leq 3 \left(\|x\|^2 + \frac{4}{\gamma^2} \|p\|^2 \right).$$

Proof The upper bound follows from

$$\| (x, p) \|^2 = \left\| x + \frac{2}{\gamma} p \right\|^2 + \|x\|^2 \leq 2 \left(\|x\|^2 + \frac{4}{\gamma^2} \|p\|^2 \right) + \|x\|^2.$$

The lower bound follows from

$$\frac{4}{\gamma^2} \|p\|^2 \leq 2 \left\| x + \frac{2}{\gamma} p \right\|^2 + 2 \|x\|^2. \quad \square$$

Consequently, the contraction result in [Theorem 5.3.6](#) for $\gamma = \sqrt{2\beta}$ implies

$$\|X_t^1 - X_t^0\|^2 + \frac{2}{\beta} \|P_t^1 - P_t^0\|^2 \leq 9 \exp\left(-\frac{\alpha t}{\sqrt{2\beta}}\right) \left(\|X_0^1 - X_0^0\|^2 + \frac{2}{\beta} \|P_0^1 - P_0^0\|^2 \right).$$

Remark 5.3.8. If we compare [Theorem 5.3.6](#) with [Optimization Box 5.3.2](#), we see that we needed to choose a larger value for the friction ($\gamma \asymp \sqrt{\beta}$ rather than $\gamma \asymp \sqrt{\alpha}$) for the proof to work, and this leads to a slower exponential contraction with rate $\alpha/\sqrt{\beta}$ (instead of $\sqrt{\alpha}$). Thus, [Theorem 5.3.6](#) can be considered an *unaccelerated* convergence rate.

5.3.2 Wasserstein Coupling Argument

We now discretize the underdamped Langevin diffusion. Of course, we could apply a simple Euler–Maruyama discretization to the SDE, but there is a slightly better discretization here. We observe that if we fix the value of the gradient term at time nh , then the rest of the SDE is a linear SDE, and can be integrated exactly. Namely, consider

$$\begin{aligned} d\hat{X}_t &= \hat{P}_t dt, \\ d\hat{P}_t &= -\nabla V(\hat{X}_{nh}) dt - \gamma \hat{P}_t dt + \sqrt{2\gamma} dB_t, \end{aligned} \quad \text{for } t \in [nh, (n+1)h]. \quad (\text{ULMC})$$

Then, the solution to the SDE is given explicitly in the following lemma.

Lemma 5.3.9 (Implementation of [ULMC](#)). *Conditioned on $(\hat{X}_{nh}, \hat{P}_{nh})$, the law of the next iterate $(\hat{X}_{(n+1)h}, \hat{P}_{(n+1)h})$ is explicitly given as $\text{normal}(M_{(n+1)h}, \Sigma)$ where*

$$M_{(n+1)h} = \begin{bmatrix} \hat{X}_{nh} + \gamma^{-1} (1 - \exp(-\gamma h)) \hat{P}_{nh} - \gamma^{-1} (h - \gamma^{-1} (1 - \exp(-\gamma h))) \nabla V(\hat{X}_{nh}) \\ \hat{P}_{nh} \exp(-\gamma h) - \gamma^{-1} (1 - \exp(-\gamma h)) \nabla V(\hat{X}_{nh}) \end{bmatrix}$$

and

$$\Sigma = \begin{bmatrix} \frac{2}{\gamma} \{ h - \frac{2}{\gamma} (1 - \exp(-\gamma h)) + \frac{1}{2\gamma} (1 - \exp(-2\gamma h)) \} & * \\ \frac{1}{\gamma} \{ 1 - 2 \exp(-\gamma h) + \exp(-2\gamma h) \} & 1 - \exp(-2\gamma h) \end{bmatrix} \otimes I_d.$$

The $*$ indicates that the entry is determined by symmetry.

The lemma is an exercise in stochastic calculus (Exercise 5.8). The point is that the discretization given by **ULMC** is implementable.

We now proceed to a discretization analysis which is a refinement of Cheng et al. (2018) that matches the guarantee of Dalalyan and Riou-Durand (2020).

Theorem 5.3.10 (Coupling analysis of **ULMC**). *For $n \in \mathbb{N}$, let $\hat{\mu}_{nh}$ denote the law of the n -th iterate of **ULMC** with step size $h \asymp \varepsilon/\sqrt{\beta\kappa d}$, friction parameter $\gamma = \sqrt{2\beta}$, and initialized at $\hat{\mu}_0 \otimes \text{normal}(0, I_d)$. Also, let $\hat{\mu}_{nh}$ denote the law of $\hat{X}_{nh}$. Assume that the target $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Then, we obtain the guarantee $\sqrt{\alpha} W_2(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$ after*

$$N = O\left(\frac{\kappa^{3/2} d^{1/2}}{\varepsilon} \log \frac{\sqrt{\alpha} W_2(\hat{\mu}_0, \pi)}{\varepsilon}\right) \text{ iterations.}$$

Proof We apply Remark 5.1.5, in the norm $\|\cdot\|$. Let $(\hat{X}_t, \hat{P}_t)_{t \geq 0}$ denote **ULMC** and let $(X_t, P_t)_{t \geq 0}$ denote the continuous-time underdamped Langevin diffusion, both driven by the same Brownian motion and started at (x, p) . We want to bound the distance $\mathbb{E}[\|(\hat{X}_h, \hat{P}_h) - (X_h, P_h)\|^2]$. According to Lemma 5.3.7, it suffices to bound $\mathbb{E}[\|\hat{X}_h - X_h\|^2]$ and $\mathbb{E}[\|\hat{P}_h - P_h\|^2]$ separately.

First,

$$\mathbb{E}[\|\hat{X}_h - X_h\|^2] = \mathbb{E}\left[\left\|\int_0^h \{\hat{P}_t - P_t\} dt\right\|^2\right] \leq h \int_0^h \mathbb{E}[\|\hat{P}_t - P_t\|^2] dt.$$

Next,

$$\begin{aligned} \mathbb{E}[\|\hat{P}_t - P_t\|^2] &= \mathbb{E}\left[\left\|\int_0^t \{-\nabla V(x) + \nabla V(X_s) - \gamma(\hat{P}_s - P_s)\} ds\right\|^2\right] \\ &\lesssim h \int_0^t (\mathbb{E}[\|\nabla V(X_s) - \nabla V(x)\|^2] + \gamma^2 \mathbb{E}[\|\hat{P}_s - P_s\|^2]) ds. \end{aligned}$$

By Grönwall's inequality, if $h \leq 1/\gamma = 1/\sqrt{2\beta}$, then

$$\mathbb{E}[\|\hat{P}_t - P_t\|^2] \lesssim h \int_0^t \mathbb{E}[\|\nabla V(X_s) - \nabla V(x)\|^2] ds \leq \beta^2 h \int_0^t \mathbb{E}[\|X_s - x\|^2] ds.$$

Again, we need a movement bound for the underdamped Langevin diffusion, which is done in Lemma 5.3.11. Substituting this in and assuming $h \ll 1/\beta^{1/2}$,

$$\begin{aligned} \mathbb{E}[\|(\hat{X}_h, \hat{P}_h) - (X_h, P_h)\|^2] &\leq \mathbb{E}[\|\hat{X}_h - X_h\|^2] + \frac{1}{\beta} \mathbb{E}[\|\hat{P}_h - P_h\|^2] \\ &\lesssim \beta h^4 \|p\|^2 + \beta^{3/2} d h^5 + \beta h^6 \|\nabla V(x)\|^2. \end{aligned}$$

This is a bound on $\mathcal{E}_{\text{strong}}(x, p)$. We apply Remark 5.1.5 with $L = \exp(-\alpha h/\sqrt{8\beta})$ (by Theorem 5.3.6) and $L_{\text{strong}} = O(\beta h^2)$, so that $L' \leq \exp(-\alpha h/4\sqrt{\beta})$ for $h \ll 1/\beta^{1/2}$. Then, since $\|\mathcal{E}_{\text{strong}}\|_{L^2(\pi)} \lesssim \beta^{1/2} d^{1/2} h^2$, we obtain

$$\mathcal{W}_2(\hat{\mu}_{Nh}, \pi) \leq \exp\left(-\frac{\alpha N h}{4\sqrt{\beta}}\right) \mathcal{W}_2(\hat{\mu}_0, \pi) + O\left(\frac{\beta d^{1/2} h}{\alpha}\right)$$

where $\mathcal{W}_2$ denotes the 2-Wasserstein distance in the $\|\cdot\|$ norm. Choosing the step size appropriately completes the proof. $\square$

The next lemma provides the movement bound (Exercise 5.9).

Lemma 5.3.11 (Underdamped movement bound). *Let $(X_t, P_t)_{t \geq 0}$ denote the underdamped Langevin diffusion with potential V satisfying $-\beta I_d \leq \nabla^2 V \leq \beta I_d$. If $t \leq 1/\sqrt{\beta} \wedge 1/\gamma$, then*

$$\mathbb{E}[\|X_t - X_0\|^2] \lesssim t^2 \mathbb{E}[\|P_0\|^2] + \gamma dt^3 + t^4 \mathbb{E}[\|\nabla V(X_0)\|^2].$$

5.3.3 Randomized Midpoint Discretization

The randomized midpoint method of Section 5.1 can be applied to the underdamped Langevin diffusion to yield an even better sampling guarantee. This was carried out in Shen and Lee (2019); Yu et al. (2024), but we state a modified variant from Altschuler et al. (2025). We set, for $t \in [0, h]$,

$$\begin{aligned} \hat{X}_{nh+t}^+ &:= \hat{X}_{nh} + \frac{1 - \exp(-\gamma t)}{\gamma} \hat{P}_{nh} - \frac{1}{\gamma} \left(t - \frac{1 - \exp(-\gamma t)}{\gamma} \right) \nabla V(\hat{X}_{nh}) + \xi_{nh,t}^{(1)}, \\ \hat{X}_{(n+1)h} &:= \hat{X}_{nh} + \frac{1 - \exp(-\gamma h)}{\gamma} \hat{P}_{nh} - \frac{1}{\gamma} \left(h - \frac{1 - \exp(-\gamma h)}{\gamma} \right) \nabla V(\hat{X}_{nh+u_{nh}}^+) + \xi_{nh,h}^{(1)}, \quad (\text{RM-ULMC}) \\ \hat{P}_{(n+1)h} &:= \exp(-\gamma h) \hat{P}_{nh} - \frac{1}{\gamma} (1 - \exp(-\gamma h)) \nabla V(\hat{X}_{nh+v_{nh}}^+) + \xi_{nh,h}^{(2)}, \end{aligned}$$

where $(u_n, v_n)_{n \in \mathbb{N}}$ is a sequence of i.i.d. random pairs, independent of the Brownian motion, with the following marginal densities on $[0, 1]$: $u_n \sim q_u$ and $v_n \sim q_v$, where

$$q_u(u) = \frac{h(1 - \exp(-\gamma(1-u)h))}{h - (1 - \exp(-\gamma h))/\gamma}, \quad q_v(v) = \frac{\gamma h \exp(-\gamma(1-v)h)}{1 - \exp(-\gamma h)}.$$

Also, we define the following Brownian integrals:

$$\begin{aligned} \xi_{nh,t}^{(1)} &:= \sqrt{\frac{2}{\gamma}} \int_{nh}^{nh+t} (1 - \exp(-\gamma(nh+t-s))) dB_s, \\ \xi_{nh,t}^{(2)} &:= \sqrt{2\gamma} \int_{nh}^{nh+t} \exp(-\gamma(nh+t-s)) dB_s. \end{aligned}$$

Although the form of the algorithm seems complicated, it is designed to facilitate the computation of the local errors. Since the proof of the following theorem is a combination of the proof techniques introduced for the randomized midpoint discretization (Theorem 5.1.1) and the analysis of underdamped Langevin (Theorem 5.3.10), we leave it as an exercise.

Theorem 5.3.12 (RM-ULMC). *For $n \in \mathbb{N}$, let $\hat{\mu}_{nh}$ denote the law of $\hat{X}_{nh}$ along RM-ULMC with step size $h > 0$. Assume that the target $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and $\nabla V(0) = 0$. If we choose $h \asymp \varepsilon^{2/3}/\beta^{1/2}d^{1/3}$, then for all sufficiently small $\varepsilon > 0$ and with an appropriate initialization, we obtain the guarantee $\sqrt{\alpha} W_2(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$ after*

$$N = \tilde{O}\left(\frac{\kappa d^{1/3}}{\varepsilon^{2/3}}\right) \quad \text{iterations}.$$

See the Bibliographical Notes for further developments.

5.3.4 Hypocoercivity

Our study of the underdamped Langevin diffusion has produced samplers with substantially improved sampling complexity, culminating in [Theorem 5.3.12](#), but we have not addressed the motivating question in this section. Namely, is there an *acceleration* phenomenon for sampling, in the sense of an algorithm with a square root dependence on the condition number? (See [Optimization Box 5.3.2](#).)

Toward this question and others, such as whether we can obtain KL divergence guarantees under functional inequality assumptions (as we did in [Section 4.2](#)), we must now enter into a discussion of the hypocoercivity arguments which until this point we had deferred due to the higher level of technical sophistication.

We start by noting that the generator of the underdamped Langevin diffusion can be written (see [Exercise 5.13](#)) as $\mathcal{L} = \gamma \mathcal{L}_{\text{OU}} + \mathcal{L}_{\text{Ham}}$, where $\mathcal{L}_{\text{OU}}$ is the generator of the Ornstein–Uhlenbeck process ([Exercise 1.3](#)) acting on the momentum coordinate,

$$\mathcal{L}_{\text{OU}} f := \Delta_p f - \langle p, \nabla_p f \rangle,$$

and $\mathcal{L}_{\text{Ham}}$ captures the Hamiltonian part of the dynamics,

$$\mathcal{L}_{\text{Ham}} f := \langle p, \nabla_x f \rangle - \langle \nabla V, \nabla_p f \rangle.$$

As it turns out, $\mathcal{L}_{\text{Ham}}$ is *anti-symmetric* in $L^2(\pi)$, i.e., $\mathcal{L}_{\text{Ham}}^* = -\mathcal{L}_{\text{Ham}}$ ([Exercise 5.14](#)). Therefore, if we write the Fokker–Planck equation for the underdamped Langevin diffusion in terms of the relative density $\rho_t := \frac{d\pi_t}{d\pi}$, we obtain the differential equation

$$\partial_t \rho_t = \mathcal{L}^* \rho_t, \quad \mathcal{L}^* = \gamma \mathcal{L}_{\text{OU}} - \mathcal{L}_{\text{Ham}}. \quad (5.3.13)$$

Now suppose that we wish to prove convergence in the chi-squared divergence. The first approach to try is to simply differentiate $t \mapsto \|\rho_t - 1\|_{L^2(\pi)}^2$, as we did for the Langevin diffusion. We find that

$$\partial_t \|\rho_t - 1\|_{L^2(\pi)}^2 = 2\gamma \langle \rho_t - 1, \mathcal{L}_{\text{OU}}(\rho_t - 1) \rangle_{L^2(\pi)} - 2 \langle \rho_t - 1, \mathcal{L}_{\text{Ham}}(\rho_t - 1) \rangle_{L^2(\pi)}.$$

Unfortunately, this naïve calculation cannot succeed. Due to the anti-symmetry of $\mathcal{L}_{\text{Ham}}$, the second term vanishes. Moreover, the first term is not *coercive*, in the sense that the putative inequality $\langle \rho_t - 1, \mathcal{L}_{\text{OU}}(\rho_t - 1) \rangle_{L^2(\pi)} \lesssim -\|\rho_t - 1\|_{L^2(\pi)}^2$ generally fails, because $\mathcal{L}_{\text{OU}}$ has a large kernel: namely, all test functions which depend only on the position coordinate.

However, there is a certain intuition that the underdamped Langevin diffusion ought to converge anyway: the $\mathcal{L}_{\text{OU}}$ operator is coercive when restricted to the momentum coordinate, and the Hamiltonian dynamics encoded in $\mathcal{L}_{\text{Ham}}$ causes the position and momentum coordinates to “mix”. Indeed, it is the *interaction* of the two parts of the dynamics that drives the system to equilibrium.

We now give an exposition of the theory of **hypocoercivity**, as treated in Villani’s monograph ([Villani, 2009a](#)), which places the above intuition into a mathematical framework. The theory has two key ingredients. First, the “interaction” between the two parts of the dynamics is handled by studying the *commutator* of the two operators. Second, in order to take advantage of the non-zero commutator, we perform an analysis in a *modified norm* (similar in spirit to the twisted metric in the coupling argument of [Theorem 5.3.6](#)). We will establish the following theorem.

Theorem 5.3.14 (Hypocoercivity). *For $t > 0$, let $f_t := \rho_t - 1 := \frac{d\rho_t}{dt} - 1$ denote the centered relative density along the underdamped Langevin diffusion with friction coefficient $\gamma = c_0\sqrt{\beta}$, $c_0 > 0$. Assume that measure $\pi \propto \exp(-V)$ on $\mathbb{R}^d$ satisfies a Poincaré inequality with constant α^{-1} , and that ∇V is β -Lipschitz. Then, there is a constant $c' > 0$ depending only on c_0 such that*

$$\begin{aligned} & \|f_t\|_{L^2(\pi)}^2 + \frac{1}{\beta} \|\nabla_x f_t\|_{L^2(\pi)}^2 + \|\nabla_p f_t\|_{L^2(\pi)}^2 \\ & \lesssim \exp\left(-\frac{c'\alpha t}{\sqrt{\beta}}\right) (\|f_0\|_{L^2(\pi)}^2 + \frac{1}{\beta} \|\nabla_x f_0\|_{L^2(\pi)}^2 + \|\nabla_p f_0\|_{L^2(\pi)}^2). \end{aligned}$$

To better appreciate the core of the argument, we follow Villani (2009a) and switch to more abstract notation. Although we are only concerned here with the underdamped Langevin diffusion, some of the following discussion will be phrased more generally to emphasize its broader applicability. We observe that $\mathcal{L}_{OU} = -\nabla_p^* \nabla_p$, which makes apparent its symmetry and non-positivity, and hence study an operator $\mathcal{L}$ of the form $\mathcal{L} = -\gamma A^* A + B$, where A and B are linear operators and B is anti-symmetric. In our application to the underdamped Langevin diffusion, this actually corresponds to the adjoint $\mathcal{L}^*$ of the generator, with

$$A := \nabla_p, \quad B := -\mathcal{L}_{\text{Ham}}, \quad C := [A, B].$$

Here, $[A, B] := AB - BA$ is the commutator. As a preliminary step, we must compute these commutators, which we leave as an important exercise (Exercise 5.14):

$$[A, A^*] = \text{id}, \quad C = -\nabla_x, \quad [B, C] = \nabla^2 V \nabla_p = \nabla^2 V A. \quad (\text{Com})$$

Here, $[A, A^*]$ is interpreted as $([\partial_{p_1}, \nabla_{p_1}^*], \dots, [\partial_{p_d}, \nabla_{p_d}^*])$. Also, C commutes with A and A^* . Finally, we let $\|\cdot\|$ denote a Hilbert norm, which is taken to be $\|\cdot\|_{L^2(\pi)}$ for underdamped Langevin.

Already, the commutator calculation proves to be interesting. Indeed, due to the anti-symmetry of B , the computation of the time derivative along the dynamics³ induced by $\mathcal{L}$ yields $2\langle f, \mathcal{L}f \rangle = -2\gamma \|Af\|^2$, and the problem is that the quantity $\|Af\|$ is not coercive, in the sense that it does not control $\|f\|$ from above. In other words, A^*A may not be strictly positive definite. For underdamped Langevin, this corresponds to $\|Af\| = \|\nabla_p f\|_{L^2(\pi)}$ and so we are missing a term—namely, the norm of the position gradient. But this position gradient shows up precisely as the commutator C .

Motivated by this observation, the plan is to instead consider $\|f\|^2 + \|Af\|^2 + \|Cf\|^2$, which corresponds to the usual squared Sobolev norm $\|f\|_{L^2(\pi)}^2 + \|\nabla f\|_{L^2(\pi)}^2$ for underdamped Langevin. Importantly, this quantity is coercive, as a consequence of the Poincaré inequality for π (which follows from the assumed Poincaré inequality for π , together with the Poincaré inequality for the standard Gaussian). In other words, rather than requiring A^*A to be positive definite, we now only require $A^*A + C^*C$ to be positive definite. In general, this process can be repeated: if $A^*A + C^*C$ is still not positive definite, we can take another commutator $C_1 = [C, B]$ and add $C_1^*C_1$, etc., and the number of times this process must be repeated corresponds to the degeneracy of the equation. For underdamped Langevin, just the addition of C^*C is enough.

However, simply differentiating $t \mapsto \|f_t\|^2 + \|Af_t\|^2 + \|Cf_t\|^2$ is still not enough to establish convergence. The other crucial ingredient is to *change the norm*. Hence, we shall consider the

³In the abstract setting, this refers to $t \mapsto \exp(\mathcal{L}t) f_0$, i.e., the semigroup with generator $\mathcal{L}$.

following Lyapunov functional

$$\mathcal{L}_t := \|f_t\|^2 + a \|Af_t\|^2 - \frac{2b}{\sqrt{\beta}} \langle Af_t, Cf_t \rangle + \frac{c}{\beta} \|Cf_t\|^2.$$

We will choose $a, b, c > 0$ to be universal constants such that $ac > b^2$, which will ensure that the above Lyapunov functional is equivalent up to universal constants to the norm $\|f\|^2 + \|Af\|^2 + \beta^{-1} \|Cf\|^2$.

As Villani writes, the key intuition is that the effect of $\mathcal{L}_{OU} = -A^*A$ is to dissipate the terms $\|f\|^2$, $\|Af\|^2$, and $\|Cf\|^2$, whereas the effect of B is to dissipate the mixed term $-\langle Af, Cf \rangle$. To see this, suppose that M is an operator which commutes with A . Then, if $\partial_t g_t = -A^*Ag_t$,

$$\partial_t \|Mg_t\|^2 = -2 \langle Mg_t, MA^*Ag_t \rangle = -2 \langle Mg_t, A^*AMg_t \rangle = -2 \|AMg_t\|^2 \leq 0.$$

On the other hand, if $\partial_t g_t = Bg_t$, from anti-symmetry, the definition of C , and (Com),

$$\begin{aligned} -\partial_t \langle Ag_t, Cg_t \rangle &= -\langle ABg_t, Cg_t \rangle - \langle Ag_t, CBg_t \rangle \\ &= -\langle ABg_t, Cg_t \rangle - \langle Ag_t, BCg_t \rangle - \langle Ag_t, [C, B]g_t \rangle \\ &= -\langle ABg_t, Cg_t \rangle + \langle BAg_t, Cg_t \rangle + \langle Ag_t, [B, C]g_t \rangle \\ &= -\|Cg_t\|^2 + \langle Ag_t, \nabla^2 V Ag_t \rangle. \end{aligned}$$

Up to an error term, we have obtained decay in the missing “ C direction”!

We now proceed to the full calculation.

Proof of Theorem 5.3.14 If we differentiate $\mathcal{L}_t$, we obtain $-\mathcal{D}(f_t)$, where we calculate

$$\mathcal{D}(f) = 2\gamma \|Af\|^2 + \underbrace{a(\cdots)}_{\text{I}} + \underbrace{\frac{2b}{\sqrt{\beta}}(\cdots)}_{\text{II}} + \underbrace{\frac{c}{\beta}(\cdots)}_{\text{III}}.$$

Our aim is to lower bound $\mathcal{D}(f_t)$ by a multiple of $\mathcal{L}_t$. Following Villani’s notation, we denote the A - and B -contributions to terms I–III by I_A, I_B , etc. We estimate each of these six terms. Based on the above intuition, we expect I_A, II_B , and III_A to contribute terms that help us, whereas the other terms are error terms. Throughout, we repeatedly use the commutator calculations in (Com).

We start with the first term.

$$\begin{aligned} I_A &= 2a\gamma \langle Af, AA^*Af \rangle = 2a\gamma (\langle Af, A^*A^2f \rangle + \|Af\|^2) = 2a\gamma (\|A^2f\|^2 + \|Af\|^2), \\ I_B &= -2a \langle Af, ABf \rangle = -2a \langle Af, BAg_t \rangle - 2a \langle Af, Cf \rangle = -2a \langle Af, Cf \rangle \\ &\geq -2a \|Af\| \|Cf\|. \end{aligned}$$

Next, for the second term, by the above calculations,

$$\begin{aligned} II_A &= -\frac{2b\gamma}{\sqrt{\beta}} (\langle AA^*Af, Cf \rangle + \langle Af, CA^*Af \rangle) \\ &= -\frac{2b\gamma}{\sqrt{\beta}} (\langle A^*A^2f, Cf \rangle + \langle Af, Cf \rangle + \langle Af, A^*CAf \rangle) \\ &= -\frac{2b\gamma}{\sqrt{\beta}} (2 \langle A^2f, ACf \rangle + \langle Af, Cf \rangle) \geq -\frac{2b\gamma}{\sqrt{\beta}} (2 \|A^2f\| \|ACf\| + \|Af\| \|Cf\|) \\ II_B &= \frac{2b}{\sqrt{\beta}} (\|Cf\|^2 - \langle Af, \nabla^2 V Af \rangle) \geq \frac{2b}{\sqrt{\beta}} (\|Cf\|^2 - \beta \|Af\|^2). \end{aligned}$$

Finally, for the third term,

$$\begin{aligned}\text{III}_A &= \frac{2c\gamma}{\beta} \langle Cf, CA^*Af \rangle = \frac{2c\gamma}{\beta} \|ACf\|^2, \\ \text{III}_B &= -\frac{2c}{\beta} \langle Cf, CBf \rangle = -\frac{2c}{\beta} (\langle Cf, BCf \rangle + \langle Cf, [C, B]f \rangle) = \frac{2c}{\beta} \langle Cf, \nabla^2 V Af \rangle \\ &\geq -2c \|Af\| \|Cf\|.\end{aligned}$$

Let $D(f) := (\|Af\|, \frac{1}{\sqrt{\beta}} \|Cf\|, \|A^2f\|, \frac{1}{\sqrt{\beta}} \|ACf\|)$. Now assume that $\gamma = c_0\sqrt{\beta}$ for some constant $c_0 > 0$. We have established the inequality $\mathcal{D}(f) \geq \langle D(f), \Sigma D(f) \rangle$ where

$$\Sigma = \sqrt{\beta} \begin{bmatrix} 2c_0 + 2ac_0 - 2b & -2a - 2c & & \\ & 2b & & \\ & & 2ac_0 & -4bc_0 \\ & & & 2cc_0 \end{bmatrix}.$$

Now we want to show that we can choose constants $a, b, c > 0$ with $ac > b^2$, depending only on c_0 , such that the symmetric part of Σ is positive definite. In this case, since Σ is block diagonal, this just amounts to $ac > b^2$ as well as $4b(c_0 + ac_0 - b) - (a + c)^2 > 0$. If we scale a, b, c down so that $\max\{a, b, c\} \ll c_0 \wedge 1$ then the second condition is clearly satisfied. With this, we have obtained

$$\partial_t \mathcal{L}_t = -\mathcal{D}(f_t) \lesssim -\sqrt{\beta} \|Af_t\|^2 - \frac{1}{\sqrt{\beta}} \|Cf_t\|^2.$$

Applying the Poincaré inequality for π (in the norm given by $\|(x, p)\|^2 := \|x\|^2 + \|p\|^2/\alpha$),

$$\partial_t \mathcal{L}_t \lesssim -\frac{\alpha}{\sqrt{\beta}} (\|Af_t\|^2 + \frac{1}{\alpha} \|Cf_t\|^2) \leq -\frac{\alpha}{2\sqrt{\beta}} (\|f_t\|^2 + \|Af_t\|^2 + \frac{1}{\alpha} \|Cf_t\|^2) \lesssim -\frac{\alpha}{\sqrt{\beta}} \mathcal{L}_t.$$

In the last line, we used the fact that $\mathcal{L}_t$ is equivalent up to constants to the Sobolev norm $\|f\|^2 + \|Af\|^2 + \frac{1}{\beta} \|Cf\|^2$. $\square$

Remarkably, there is a variant of this theorem which holds for KL convergence, which proceeds via a twisted version of the Fisher information (see [Exercise 5.15](#) for a challenge).

Theorem 5.3.15 (Entropic hypocoercivity). *In the setting of [Theorem 5.3.14](#), suppose instead that π satisfies a log-Sobolev inequality with constant α^{-1} . Then, there is a Lyapunov function*

$$\mathcal{L}_t := \text{KL}(\pi_t \parallel \pi) + \mathbb{E}_{\pi_t} \left\langle \nabla \log \frac{\pi_t}{\pi}, \left(\begin{bmatrix} a/\beta & b/\sqrt{\beta} \\ b/\sqrt{\beta} & c \end{bmatrix} \otimes I_d \right) \nabla \log \frac{\pi_t}{\pi} \right\rangle$$

such that for a universal constant $c' > 0$ and all $t \geq 0$,

$$\mathcal{L}_t \leq \exp\left(-\frac{c'\alpha t}{\sqrt{\beta}}\right) \mathcal{L}_0.$$

Nevertheless, [Theorem 5.3.14](#) and [Theorem 5.3.15](#) only yield *unaccelerated* rates of convergence.

5.3.5 Accelerated convergence

Very recently, in a sequence of remarkable works, accelerated convergence rates were established via clever choices of Lyapunov functions. Although the following results were not the first in terms of precedence, they are somewhat more direct; see the Bibliographical Notes for a brief history.

Theorem 5.3.16 (Accelerated convergence). *Suppose that π is log-concave and satisfies mild growth/regularity conditions (to ensure that certain computations are justified), and let $\gamma = C\sqrt{\alpha}$. Then, there is a constant $c > 0$, depending only on C , such that the following hold for the underdamped Langevin diffusion.*

1 (Fan et al. (2026)) *If π satisfies a Poincaré inequality with constant α^{-1} , then for all $t \geq 0$,*

$$\chi^2(\pi_t \parallel \pi) \lesssim \exp(-c\sqrt{\alpha}t) \chi^2(\pi_0 \parallel \pi).$$

2 (Lu (2026)) *If π satisfies a log-Sobolev inequality with constant α^{-1} , then for all $t \geq 0$,*

$$\text{KL}(\pi_t \parallel \pi) \lesssim \exp(-c\sqrt{\alpha}t) \text{KL}(\pi_0 \parallel \pi).$$

Together with the result of Altschuler et al. (2025), which provided KL divergence discretization guarantees for RM-ULMC, it implies the following algorithmic result.

Theorem 5.3.17 (Altschuler et al. (2025)). *Assume that the target $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and $\nabla V(0) = 0$. Then, RM-ULMC achieves $\|\hat{\mu} - \pi\|_{\text{TV}} \leq \varepsilon$ after*

$$N = \tilde{O}\left(\frac{\kappa^{5/6} d^{1/3}}{\varepsilon^{2/3}}\right) \quad \text{iterations}.$$

The $\kappa^{5/6}$ dependence on the condition number is the current state-of-the-art progress toward acceleration for sampling; see the Bibliographical Notes for further discussion.

In the remainder of this section, we present most of the proof of Theorem 5.3.16, omitting some technical details (e.g., on regularity). The χ^2 result is based on the Lyapunov function

$$\mathcal{L}_{\chi^2}(\mu) := \frac{1}{2} \chi^2(\mu \parallel \pi) - c' \sqrt{\alpha} \left\langle \frac{\mu}{\pi}, \mathcal{A} \frac{\mu}{\pi} \right\rangle_{L^2(\pi)}, \quad \mathcal{A} := (\alpha - \mathcal{L})^{-1} \Pi \mathcal{L}_{\text{Ham}},$$

where Π is the momentum averaging operator: $\Pi f(x) := \int f(x, p) n(dp)$, with $n := \text{normal}(0, I_d)$, and $\mathcal{L}$ is the overdamped Langevin generator. On the other hand, the KL result is based on the Lyapunov function

$$\mathcal{L}_{\text{KL}}(\mu) := \text{KL}(\mu \parallel \pi) + c' \sqrt{\alpha} \int \langle p, x - T_{\mu^x \rightarrow \pi}(x) \rangle \mu(dx, dp),$$

where $T_{\mu^x \rightarrow \pi}$ is the optimal transport map from the X -marginal μ^X of μ to π . The precise relationship between these Lyapunov functions is that $\mathcal{L}_{\text{KL}}$ almost linearizes to $\mathcal{L}_{\chi^2}$; see Exercise 5.16. For brevity, we focus on the proof of the second statement.

Proof of Theorem 5.3.16, log-Sobolev case. Along the underdamped Langevin diffusion (5.3.1), we already know that $\partial_t \text{KL}(\pi_t \parallel \pi) = -\gamma \int \|\nabla_p \log(\pi_t/\pi)\|^2 d\pi_t =: -\gamma \text{FI}^P(\pi_t \parallel \pi)$, so it remains to differentiate the second corrector term.

Define the vector field $\bar{p}_t(x) := \mathbb{E}[P_t \mid X_t = x]$, where $(X_t, P_t)_{t \geq 0}$ is the underdamped Langevin diffusion. Then, $\partial_t \mathbb{E} \varphi(X_t) = \mathbb{E} \langle \nabla \varphi(X_t), P_t \rangle = \mathbb{E} \langle \nabla \varphi(X_t), \bar{p}_t(X_t) \rangle$, so by [Theorem 1.4.6](#),

$$\mathcal{C}(\pi_t) := \mathbb{E} \langle P_t, X_t - T_{\pi_t^X}(X_t) \rangle = \mathbb{E} \langle \bar{p}_t(X_t), X_t - T_{\pi_t^X}(X_t) \rangle = \frac{1}{2} \partial_t W_2^2(\pi_t^X, \pi),$$

where we abbreviate $T_{\pi_t^X} := T_{\pi_t^X \rightarrow \pi}$. Next, define the curve $(x_s)_{s \geq t}$ so that $x_t \sim \pi_t^X$ and $z \sim \pi$ are optimally coupled, and $\dot{x}_s = \bar{p}_s(x_s)$. This also implies that $\ddot{x}_s = \partial_s[\bar{p}_s(x_s)] = \partial_s \bar{p}_s(x_s) + \nabla \bar{p}_s(x_s) \bar{p}_s(x_s)$. Since $x_s \sim \pi_s^X$ for all $s \geq t$ by the continuity equation (1.3.18),

$$\begin{aligned} W_2^2(\pi_{t+h}^X, \pi) &\leq \mathbb{E}[\|x_{t+h} - z\|^2] = \mathbb{E}[\|x_t + h \bar{p}_t(x_t) + \frac{h^2}{2} [\partial_t \bar{p}_t(x_t) + \nabla \bar{p}_t(x_t) \bar{p}_t(x_t)] + o(h^2) - z\|^2] \\ &= W_2^2(\pi_t^X, \pi) + 2h \mathbb{E} \langle x_t - z, \bar{p}_t(x_t) \rangle \\ &\quad + h^2 \mathbb{E}[\|\bar{p}_t(x_t)\|^2 + \langle x_t - z, \partial_t \bar{p}_t(x_t) + \nabla \bar{p}_t(x_t) \bar{p}_t(x_t) \rangle] + o(h^2). \end{aligned}$$

Since we can identify the second term above as $h \partial_t W_2^2(\pi_t^X, \pi)$, it implies

$$\partial_t \mathcal{C}(\pi_t) = \frac{1}{2} \partial_t^2 W_2^2(\pi_t^X, \pi) \leq \mathbb{E}[\|\bar{p}_t(x_t)\|^2 + \langle x_t - z, \partial_t \bar{p}_t(x_t) + \nabla \bar{p}_t(x_t) \bar{p}_t(x_t) \rangle].$$

Next, a tedious computation using (5.3.4) and the continuity equation $\partial_t \pi_t^X + \operatorname{div}(\bar{p}_t \pi_t^X) = 0$ yields an identity for the derivative $\partial_t \bar{p}_t(x) = \partial_t [\int p \pi_t(x, p) dp / \pi_t^X(x)]$:

$$\partial_t \bar{p}_t + \nabla \bar{p}_t \bar{p}_t = -\gamma \bar{p}_t - \nabla_x \cdot \Sigma_t - \Sigma_t \nabla \log \pi_t^X - \nabla V, \quad (5.3.18)$$

where $\Sigma_t(x)$ is the conditional covariance of P_t given $X_t = x$; see [Exercise 5.17](#). Thus, by integration by parts,

$$\begin{aligned} \partial_t \mathcal{C}(\pi_t) &\leq -\gamma \mathcal{C}(\pi_t) + \int \{\|\bar{p}_t\|^2 - \langle \operatorname{id} - T_{\pi_t^X}, \nabla_x \cdot \Sigma_t + \Sigma_t \nabla \log \pi_t^X + \nabla V \rangle\} d\pi_t^X \\ &= -\gamma \mathcal{C}(\pi_t) + \int \{\|\bar{p}_t\|^2 + \langle I_d - \nabla T_{\pi_t^X}, \Sigma_t \rangle - \langle \operatorname{id} - T_{\pi_t^X}, \nabla V \rangle\} d\pi_t^X. \end{aligned}$$

We apply the following inequality pointwise: for any $m \in \mathbb{R}^d$ and $M > 0$, if μ has mean m and covariance Σ , then

$$\|m\|^2 + \langle I_d - M, \Sigma \rangle \leq 2 \operatorname{KL}(\mu \parallel \mathfrak{n}) - \log \det M; \quad (5.3.19)$$

see [Exercise 5.17](#). Applying this to $\mu = \pi_t^{P|X=x}$ and invoking the Monge–Ampère equation and the Gaussian log-Sobolev inequality yields

$$\begin{aligned} \partial_t \mathcal{C}(\pi_t) &\leq -\gamma \mathcal{C}(\pi_t) + \int \{2 \operatorname{KL}(\pi_t^{P|X=x} \parallel \mathfrak{n}) - \log \det \nabla T_{\pi_t^X}(x) - \langle x - T_{\pi_t^X}(x), \nabla V(x) \rangle\} \pi_t^X(dx) \\ &= -\gamma \mathcal{C}(\pi_t) + \int \{2 \operatorname{KL}(\pi_t^{P|X=x} \parallel \mathfrak{n}) - \log \pi_t^X(x) - V(T_{\pi_t^X}(x)) - \langle x - T_{\pi_t^X}(x), \nabla V(x) \rangle\} \pi_t^X(dx) \\ &= -\gamma \mathcal{C}(\pi_t) + \int \{2 \operatorname{KL}(\pi_t^{P|X=x} \parallel \mathfrak{n}) - V(x) - \log \pi_t^X(x) - D_V(T_{\pi_t^X}(x), x)\} \pi_t^X(dx) \\ &= -\gamma \mathcal{C}(\pi_t) + \int \{2 \operatorname{KL}(\pi_t^{P|X=x} \parallel \mathfrak{n}) - D_V(T_{\pi_t^X}(x), x)\} \pi_t^X(dx) - \operatorname{KL}(\pi_t^X \parallel \pi) \\ &\leq -\gamma \mathcal{C}(\pi_t) + \int \operatorname{Fl}(\pi_t^{P|X=x} \parallel \mathfrak{n}) \pi_t^X(dx) - \operatorname{KL}(\pi_t^X \parallel \pi) \\ &= -\gamma \mathcal{C}(\pi_t) + \operatorname{Fl}^P(\pi_t \parallel \pi) - \operatorname{KL}(\pi_t^X \parallel \pi). \end{aligned}$$

Note that $-\text{KL}(\pi_t^X \parallel \pi)$ is the missing term that we need to close the differential inequality. Indeed,

$$\partial_t \mathcal{L}_{\text{KL}}(\pi_t) \leq -(\gamma - c' \sqrt{\alpha}) \text{Fl}^P(\pi_t \parallel \pi) - c' \sqrt{\alpha} \gamma \mathcal{C}(\pi_t) - c' \sqrt{\alpha} \text{KL}(\pi_t^X \parallel \pi).$$

We have the bounds

$$\begin{aligned} |\mathcal{C}(\pi_t)| &= \left| \int \langle \bar{p}_t, \text{id} - T_{\pi_t^X} \rangle d\pi_t^X \right| \leq \|\bar{p}_t\|_{L^2(\pi_t^X)} W_2(\pi_t^X, \pi) \leq \|W_2(\pi_t^{P|X=\bullet}, n)\|_{L^2(\pi_t^X)} W_2(\pi_t^X, \pi) \\ &\leq \frac{2}{\sqrt{\alpha}} \sqrt{\text{KL}(\pi_t^X \parallel \pi) \int \text{KL}(\pi_t^{P|X=x} \parallel n) \pi_t^X(dx)} \\ &\leq \frac{1}{\sqrt{\alpha}} \min\{\text{KL}(\pi_t \parallel \pi), \sqrt{2 \text{KL}(\pi_t^X \parallel \pi) \text{Fl}^P(\pi_t \parallel \pi)}\}, \end{aligned}$$

where, in the last line, the first bound is from Young's inequality and the second from the Gaussian log-Sobolev inequality. The first inequality implies, in particular, that $1 - c' \leq \mathcal{L}_{\text{KL}}(\mu)/\text{KL}(\mu \parallel \pi) \leq 1 + c'$. Taking $c' = C/\{2(1 + C^2)\}$,

$$\begin{aligned} \partial_t \mathcal{L}_{\text{KL}}(\pi_t) &\leq -(C - c')\sqrt{\alpha} \text{Fl}^P(\pi_t \parallel \pi) - c' \sqrt{\alpha} \gamma \mathcal{C}(\pi_t) - c' \sqrt{\alpha} \text{KL}(\pi_t^X \parallel \pi) \\ &\leq -(C - c')\sqrt{\alpha} \text{Fl}^P(\pi_t \parallel \pi) + Cc' \sqrt{2\alpha \text{KL}(\pi_t^X \parallel \pi) \text{Fl}^P(\pi_t \parallel \pi)} - c' \sqrt{\alpha} \text{KL}(\pi_t^X \parallel \pi) \\ &\leq -\sqrt{\alpha} \left[(C - c' - C^2 c') \text{Fl}^P(\pi_t \parallel \pi) + \frac{c'}{2} \text{KL}(\pi_t^X \parallel \pi) \right] \\ &= -\sqrt{\alpha} \left[\frac{C}{2} \text{Fl}^P(\pi_t \parallel \pi) + \frac{c'}{2} \text{KL}(\pi_t^X \parallel \pi) \right] \leq -\frac{c' \sqrt{\alpha}}{2} \text{KL}(\pi_t \parallel \pi) \leq -\frac{c' \sqrt{\alpha}}{2(1 + c')} \mathcal{L}_{\text{KL}}(\pi_t). \end{aligned}$$

Hence,

$$\begin{aligned} \text{KL}(\pi_t \parallel \pi) &\leq \frac{1}{1 - c'} \mathcal{L}_{\text{KL}}(\pi_t) \leq \frac{1}{1 - c'} \exp\left(-\frac{c' \sqrt{\alpha}}{2(1 + c')} t\right) \mathcal{L}_{\text{KL}}(\pi_0) \\ &\leq \frac{1 + c'}{1 - c'} \exp\left(-\frac{c' \sqrt{\alpha}}{2(1 + c')} t\right) \text{KL}(\pi_0 \parallel \pi). \quad \square \end{aligned}$$

5.3.6 Second-order lifts and space-time functional inequalities

We present another approach to accelerated convergence based on space-time functional inequalities. This approach, which actually predated the proof from the preceding subsection, requires some more machinery, but is illuminating in its own right and yields more information. We follow the second-order lifting framework of [Eberle and Lörler \(2026\)](#).

Definition 5.3.20 (Second-order lift). Let $(P_t)_{t \geq 0}$ be a Markov semigroup over a space $\mathcal{X}$, with generator $\mathcal{L}$, Dirichlet form $\mathcal{E}$, and which is reversible with respect to a probability measure π . Let $(\mathbf{P}_t)_{t \geq 0}$ be a Markov semigroup over a space $\mathcal{X} \times \mathcal{P}$, with generator $\mathcal{L}$ and stationary distribution π . We denote by $\Pi : \mathcal{X} \times \mathcal{P} \rightarrow \mathcal{X}$ the projection map $(x, p) \mapsto x$.

We say that $(\mathbf{P}_t)_{t \geq 0}$ is a **second-order lift** of $(P_t)_{t \geq 0}$ if $\Pi_{\#} \pi = \pi$, and for all $f, g \in \text{dom } \mathcal{L}$,

$$\int (f \circ \Pi) \mathcal{L}(g \circ \Pi) d\pi = 0, \quad \int \mathcal{L}(f \circ \Pi) \mathcal{L}(g \circ \Pi) d\pi = \mathcal{E}(f, g).$$

Proposition 5.3.21 (Underdamped Langevin as a second-order lift). *The underdamped Langevin diffusion (5.3.1) is a second-order lift of the Langevin diffusion (1.0.1).*

Proof Here, $\mathcal{L} = \mathcal{L}_{\text{Ham}} + \gamma \mathcal{L}_{\text{OU}}$. Note that $\mathcal{L}_{\text{OU}}$ only acts on the momentum coordinate, whereas the definition of a second-order lift only cares about functions of the position variable; moreover, both $\mathcal{L}_{\text{Ham}}$ and $\mathcal{L}_{\text{OU}}$ leave π stationary. Hence, it suffices to check that $\mathcal{L}_{\text{Ham}}$ is a second-order lift of $\mathcal{L}$, where $\mathcal{L}$ is the Langevin generator.

The first condition is true because

$$\int (f \circ \Pi) \mathcal{L}_{\text{Ham}}(g \circ \Pi) d\pi = \int (f \circ \Pi) \langle p, \nabla_x g \rangle d\pi = 0,$$

by integrating over p first. For the second condition,

$$\int \mathcal{L}_{\text{Ham}}(f \circ \Pi) \mathcal{L}_{\text{Ham}}(g \circ \Pi) d\pi = \int \langle p, \nabla_x f \rangle \langle p, \nabla_x g \rangle d\pi = \int \langle \nabla_x f, \nabla_x g \rangle d\pi,$$

which is the Dirichlet form for the Langevin diffusion. $\square$

The general idea here is that starting with a reversible Markov process, we can lift it to a new Markov process, which is usually non-reversible and can speed up convergence. Interestingly, we can prove a lower bound on the relaxation time of *any* second-order lift, showing that a square-root speed up is the best possible.

Theorem 5.3.22 (Relaxation time lower bound for second-order lifts, [Eberle and Lörler \(2026\)](#)).

Let $(P_t)_{t \geq 0}$ be a second-order lift of $(P_t)_{t \geq 0}$. Define the relaxation time of $\mathcal{L}$ to be

$$\text{relax}(\mathcal{L}) := \inf \left\{ t \geq 0 \mid \|P_t f\|_{L^2(\pi)} \leq e^{-1} \|f\|_{L^2(\pi)} \text{ for all } f \in L^2(\pi) \text{ with } \int f d\pi = 0 \right\},$$

and similarly define the relaxation time of $\mathcal{L}$. Then,

$$\text{relax}(\mathcal{L}) \gtrsim \sqrt{\text{relax}(\mathcal{L})}.$$

See [Exercise 5.18](#). Combined with [Theorem 5.3.16](#), this justifies that the underdamped Langevin diffusion is an *optimal* second-order lift of the Langevin diffusion.

We now proceed to outline a proof of the Poincaré case of [Theorem 5.3.16](#). From now on, we specialize to the setting in which $\mathcal{L}$ is the generator of the underdamped Langevin diffusion, and write $\pi = \pi \otimes n$ where $n := \text{normal}(0, I_d)$. Let m_T denote the Lebesgue measure over $[0, T]$, and for $f : \mathbb{R}_+ \times \mathbb{R}^d \rightarrow \mathbb{R}$ write

$$\mathcal{E}_T(f) := \int_0^T \mathcal{E}(f(t, \cdot)) dt = \int_0^T \|\nabla_x f(t, \cdot)\|_{L^2(\pi)}^2 dt.$$

Similarly, for $f : \mathbb{R}_+ \times \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$, write

$$\mathcal{E}_T(f) := \int_0^T \mathcal{E}(f(t, \cdot)) dt = \gamma \int_0^T \|\nabla_p f(t, \cdot)\|_{L^2(\pi)}^2 dt.$$

For a function $g : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$, we let Π denote the momentum averaging operator, $\Pi g(x) := \int g(x, p) n(dp)$. We abuse notation by adding or removing arguments from functions as necessary: for instance, if $g : \mathbb{R}_+ \times \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$, we can apply the momentum averaging operator to obtain

$\Pi g(t, x) := \int g(t, x, p) n(dp)$; and we can further think of Πg as a function defined on $\mathbb{R}_+ \times \mathbb{R}^d \times \mathbb{R}^d$ via $\Pi g(t, x, p) := \Pi g(t, x)$, i.e., it is constant with respect to the momentum variable.

We begin with preliminary calculations.

Lemma 5.3.23 (Preliminary calculations).

1 Assume that π is log-concave. For functions $f : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$ and $g : \mathbb{R}^d \rightarrow \mathbb{R}$ (depending on the position coordinate only),

$$|\langle \mathcal{L}_{\text{Ham}}(f - \Pi f), \mathcal{L}_{\text{Ham}} g \rangle_{L^2(\pi)}| \leq \frac{2}{\sqrt{\gamma}} \sqrt{\mathcal{E}(f)} \|\mathcal{L} g\|_{L^2(\pi)}.$$

2 (**Divergence lemma**) Assume that π satisfies a Poincaré inequality with constant $1/\alpha$. For any $f \in L^2(\mathfrak{m}_T \otimes \pi)$ with $\int f d(\mathfrak{m}_T \otimes \pi) = 0$, there exist functions g, h over $[0, T] \times \mathbb{R}^d$ with $g(0, \cdot) = g(T, \cdot) = h(0, \cdot) = h(T, \cdot) = 0$ such that $f = \partial_t g + \mathcal{L} h$, satisfying the estimates

$$\sqrt{\mathcal{E}_T(g)} + \|\mathcal{L} h\|_{L^2(\mathfrak{m}_T \otimes \pi)} \lesssim \|f\|_{L^2(\mathfrak{m}_T \otimes \pi)}, \quad \sqrt{\mathcal{E}_T(\partial_t h)} \lesssim \left(1 + \frac{1}{\sqrt{\alpha} T}\right) \|f\|_{L^2(\mathfrak{m}_T \otimes \pi)}.$$

Proof

1 The Gaussian Poincaré inequality implies

$$\begin{aligned} |\langle \mathcal{L}_{\text{Ham}}(f - \Pi f), \mathcal{L}_{\text{Ham}} g \rangle_{L^2(\pi)}| &= |\langle f - \Pi f, (\text{id} - \Pi) \mathcal{L}_{\text{Ham}}^2 g \rangle_{L^2(\pi)}| \\ &\leq \|f - \Pi f\|_{L^2(\pi)} \|(\text{id} - \Pi) \mathcal{L}_{\text{Ham}}^2 g\|_{L^2(\pi)} \\ &\leq \sqrt{\frac{\mathcal{E}(f)}{\gamma}} \|(\text{id} - \Pi) \mathcal{L}_{\text{Ham}}^2 g\|_{L^2(\pi)}. \end{aligned}$$

A computation shows that $\mathcal{L}_{\text{Ham}}^2 g = \langle p, \nabla_x^2 g p \rangle - \langle \nabla V, \nabla_x g \rangle$, so that

$$\begin{aligned} \|(\text{id} - \Pi) \mathcal{L}_{\text{Ham}}^2 g\|_{L^2(\pi)}^2 &= \int \text{var}_\pi(\langle p, \nabla_x^2 g p \rangle - \langle \nabla V, \nabla_x g \rangle) d\pi \leq 4 \int \|\nabla_x^2 g p\|^2 d\pi \\ &= 4 \int \|\nabla_x^2 g\|_{\text{HS}}^2 d\pi = 4 \int \{\Gamma_2(g) - \langle \nabla_x g, \nabla^2 V \nabla_x g \rangle\} d\pi \leq 4 \|\mathcal{L} g\|_{L^2(\pi)}^2. \end{aligned}$$

2 The divergence lemma is a technical result in PDE theory which we do not prove here, but we discuss some considerations. A first approach would be to solve the PDE $\overline{\mathcal{L}} u = -(\partial_t^2 + \mathcal{L})u = f$ on $[0, T] \times \mathbb{R}^d$ with Neumann boundary conditions $\partial_t u(0, \cdot) = \partial_t u(T, \cdot) = 0$, and then take $g := -\partial_t u$, $h := -u$. With these choices, g satisfies the correct boundary conditions, but h does not. Nevertheless, if this approach were valid, then we could establish the desired estimates using standard arguments. Recall from (2.2.5) that

$$\int \|\nabla^2 h\|_{\text{HS}}^2 d\pi = \int \{\Gamma_2(h) - \langle \nabla h, \nabla^2 V \nabla h \rangle\} d\pi = \int \{(\mathcal{L} h)^2 - \langle \nabla h, \nabla^2 V \nabla h \rangle\} d\pi. \quad (5.3.24)$$

For the operator $\overline{\mathcal{L}}$, we have a similar formula

$$\int \|\overline{\nabla}^2 u\|_{\text{HS}}^2 d(\mathfrak{m}_T \otimes \pi) = \int \{(\overline{\mathcal{L}} u)^2 - \langle \nabla_x u, \nabla^2 V \nabla_x u \rangle\} d(\mathfrak{m}_T \otimes \pi) + \text{non-pos. term},$$

where $\overline{\nabla} = (\partial_t, \nabla_x)$ denotes the gradient operator on $[0, T] \times \mathbb{R}^d$, and the last term, non-positive,

arises due to the boundary. If we suppose, for the sake of argument, that $\nabla^2 V \geq 0$, then we have $\int \|\bar{\nabla}^2 u\|_{\text{HS}}^2 d(\mathfrak{m}_T \otimes \pi) \leq \|f\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2$, which yields various estimates. For example,

$$\mathcal{E}_T(g) = \mathcal{E}_T(\partial_t h) = \|\nabla_x \partial_t u\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2,$$

so these terms would be controlled by $\|\bar{\nabla}^2 u\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2$, and $\|\mathcal{L}h\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2$ would be controlled with the help of (5.3.24).

So, the issue is that the choices above do not fully satisfy the boundary conditions. Observe that if we also had $h(0, \cdot) = h(T, \cdot) = 0$, it would imply that f is orthogonal to the space $\mathbb{H}_0$ of functions v satisfying the PDE $(\partial_t^2 + \mathcal{L})v = 0$. This motivates performing an orthogonal decomposition of f into the part that belongs to $\mathbb{H}_0$, and the part that belongs to $\mathbb{H}_0^\perp$. We refer to Eberle and Lörler (2026) for a full proof.

□

Proof of Theorem 5.3.16, Poincaré case. Our goal is to prove the *time-averaged* decay

$$\int_0^T \|\mathbf{P}_t^* f_0\|_{L^2(\pi)}^2 dt \leq C_T \int_0^T \mathcal{E}(\mathbf{P}_t^* f_0, \mathbf{P}_t^* f_0) dt = \gamma C_T \int_0^T \|\nabla_p \mathbf{P}_t^* f_0\|_{L^2(\pi)}^2 dt, \quad (5.3.25)$$

for a suitable choice of $T > 0$ and for all f_0 with $\int f_0 d\pi = 0$. In turn, this implies (Exercise 5.19)

$$\chi^2(\pi_t \parallel \pi) \leq \exp\left(-\frac{2(t-T)}{C_T}\right) \chi^2(\pi_0 \parallel \pi). \quad (5.3.26)$$

Write $f(t, x, p) := \mathbf{P}_t^* f_0(x, p)$. We start with the decomposition

$$\|f\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2 = \|\Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2 + \|f - \Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2$$

and we control each term in turn. The second term is easier, as we can apply the Poincaré inequality for the Gaussian measure π :

$$\|f - \Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2 \leq \|\nabla_p f\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2 = \frac{\mathcal{E}_T(f)}{\gamma}.$$

For the first term, we apply the divergence lemma (Lemma 5.3.23) to write $-\Pi f = \partial_t g + \mathcal{L}h$ for suitable functions g, h . We can freely integrate by parts without boundary terms due to the properties of g and h , so

$$\begin{aligned} \|\Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)}^2 &= -\langle \Pi f, \partial_t g + \mathcal{L}h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} = \langle \partial_t \Pi f, g \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} + \mathcal{E}_T(\Pi f, h) \\ &= \langle (\partial_t + \mathcal{L}_{\text{Ham}}) \Pi f, g + \mathcal{L}_{\text{Ham}} h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)}, \end{aligned}$$

where the last line follows from the second-order lift property.

We further write $\Pi f = f + \Pi f - f$ and bound terms separately. Using the dynamic equation for f , $\partial_t f = (-\mathcal{L}_{\text{Ham}} + \gamma \mathcal{L}_{\text{OU}})f$, and $\langle \mathcal{L}_{\text{OU}} f, g \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} = 0$ since g does not depend on p . So,

$$\langle (\partial_t + \mathcal{L}_{\text{Ham}}) f, g + \mathcal{L}_{\text{Ham}} h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} = \gamma \langle \mathcal{L}_{\text{OU}} f, g + \mathcal{L}_{\text{Ham}} h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} = \gamma \langle \mathcal{L}_{\text{OU}} f, \mathcal{L}_{\text{Ham}} h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)}$$

and

$$\begin{aligned}
\langle \mathcal{L}_{\text{OU}} f, \mathcal{L}_{\text{Ham}} h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} &= -\langle \nabla_p^* \nabla_p f, p \cdot \nabla_x h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} = -\langle \nabla_p f, \nabla_x h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} \\
&\leq \|\nabla_p f\|_{L^2(\mathfrak{m}_T \otimes \pi)} \|\nabla_x h\|_{L^2(\mathfrak{m}_T \otimes \pi)} = \sqrt{\frac{\mathcal{E}_T(f)}{\gamma}} \int_0^T \langle h, (-\mathcal{L}) h \rangle_{L^2(\pi)} \\
&\leq \sqrt{\frac{\mathcal{E}_T(f)}{\alpha \gamma}} \|\mathcal{L} h\|_{L^2(\mathfrak{m}_T \otimes \pi)} \lesssim \sqrt{\frac{\mathcal{E}_T(f)}{\alpha \gamma}} \|\Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)}. \quad (5.3.27)
\end{aligned}$$

For the other term, we apply Lemma 5.3.23 to obtain

$$\begin{aligned}
&| \langle (\partial_t + \mathcal{L}_{\text{Ham}})(\Pi f - f), g + \mathcal{L}_{\text{Ham}} h \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} | \\
&\leq | \langle f - \Pi f, \partial_t g + \mathcal{L}_{\text{Ham}} \partial_t h + \mathcal{L}_{\text{Ham}} g \rangle_{L^2(\mathfrak{m}_T \otimes \pi)} | + \frac{2}{\sqrt{\gamma}} \sqrt{\mathcal{E}_T(f)} \|\mathcal{L} h\|_{L^2(\mathfrak{m}_T \otimes \pi)} \\
&\lesssim \|f - \Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)} (\|\mathcal{L}_{\text{Ham}} \partial_t h\|_{L^2(\mathfrak{m}_T \otimes \pi)} + \|\mathcal{L}_{\text{Ham}} g\|_{L^2(\mathfrak{m}_T \otimes \pi)}) + \sqrt{\frac{\mathcal{E}_T(f)}{\gamma}} \|\Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)} \\
&\leq \sqrt{\frac{\mathcal{E}_T(f)}{\gamma}} (\sqrt{\mathcal{E}_T(g)} + \sqrt{\mathcal{E}_T(\partial_t h)} + \|\Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)}) \lesssim (1 + \frac{1}{\sqrt{\alpha} T}) \sqrt{\frac{\mathcal{E}_T(f)}{\gamma}} \|\Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)}.
\end{aligned}$$

We have shown that

$$\|\Pi f\|_{L^2(\mathfrak{m}_T \otimes \pi)} \lesssim \frac{1}{\sqrt{\gamma}} \left(1 + \frac{\gamma}{\sqrt{\alpha}} + \frac{1}{\sqrt{\alpha} T}\right) \sqrt{\mathcal{E}_T(f)}$$

and hence

$$\|f\|_{L^2(\mathfrak{m}_T \otimes \pi)} \lesssim \frac{1}{\sqrt{\gamma}} \left(1 + \frac{\gamma}{\sqrt{\alpha}} + \frac{1}{\sqrt{\alpha} T}\right) \sqrt{\mathcal{E}_T(f)}.$$

If we choose $\gamma \asymp 1/T \asymp \sqrt{\alpha}$, this verifies (5.3.25) with $C_T \asymp 1/\sqrt{\alpha}$. $\square$

It is natural to ask for a log-Sobolev analogue of the space-time approach, and this was established recently in Li and Lu (2026). The key is the following inequality along *controlled* paths.

Theorem 5.3.28 (Space-time log-Sobolev inequality, Li and Lu (2026)). *Let $(\mu_t)_{t \geq 0}$ be a sufficiently regular curve of probability measures such that for $\rho_t := \frac{d\mu_t}{d\pi}$ and $\pi := \pi \otimes \mathfrak{n}$,*

$$(\partial_t + \mathcal{L}_{\text{Ham}}) \rho_t = \nabla_p^* (\rho_t u_t) \quad (5.3.29)$$

for vector fields $(u_t)_{t \geq 0}$. Let π satisfy a log-Sobolev inequality with constant α^{-1} . Then, for all $T \gg 1/\sqrt{\alpha}$, it holds that

$$\frac{1}{T} \int_0^T \text{KL}(\mu_t \parallel \pi) dt \lesssim \frac{1}{T} \int_0^T \text{Fl}^P(\mu_t \parallel \pi) dt + \frac{1}{\alpha T} \int_0^T \mathbb{E}_{\mu_t} [\|u_t\|^2] dt.$$

The proof is a modification of the proof in Section 5.3.5; see Exercise 5.22. To see the relationship with the space-time Poincaré proof, note that along the underdamped Langevin diffusion, by (5.3.13),

we can take $u_t = -\gamma \nabla_p \log \rho_t$ for $\gamma \asymp \sqrt{\alpha}$ and prove the time-averaged inequality

$$\frac{1}{T} \int_0^T \text{KL}(\pi_t \parallel \pi) dt \lesssim \frac{1}{T} \int_0^T \text{FI}^P(\pi_t \parallel \pi) dt.$$

This is precisely the log-Sobolev analogue of (5.3.25), and readily implies the second statement of Theorem 5.3.16.

However, as shown by Li and Lu (2026), Theorem 5.3.28 yields more: it also implies accelerated hypercontractivity. When we translate this result in terms of Rényi divergences, it yields the following result; see Exercise 5.23.

Theorem 5.3.30 (Accelerated hypercontractivity, Li and Lu (2026)). *Suppose that π is log-concave, satisfies a log-Sobolev inequality with constant α^{-1} , and satisfies mild growth/regularity conditions. Let $\gamma = C\sqrt{\alpha}$. Then, there are constants $c, \bar{C} > 0$ depending only on C such that, for all $q > 1$, along the underdamped Langevin diffusion,*

$$\mathcal{R}_q(\pi_t \parallel \pi) \leq \bar{C} q \exp(-c\sqrt{\alpha} t) \mathcal{R}_q(\pi_0 \parallel \pi).$$

Moreover, if $q - 1 < q' - 1 \leq \exp(\bar{C}\sqrt{\alpha} t) (q - 1)$, and $t \gg 1/\sqrt{\alpha}$, then

$$\mathcal{R}_{q'}(\pi_t \parallel \pi) \leq \frac{q'(q-1)}{q(q'-1)} \mathcal{R}_q(\pi_0 \parallel \pi).$$

Bibliographical Notes

The local error framework is classical (Milstein and Tretyakov, 2021) and has been widely used for the analysis of sampling algorithms (Erdogdu et al., 2018; Li et al., 2019; Sanz-Serna and Zygalakis, 2021; Li et al., 2022a,b). The improved complexity for LMC in Exercise 5.2 was obtained in Li et al. (2022b), whereas the rate when $\nabla^2 V$ is Lipschitz in the operator norm goes back to earlier works.

HMC and its variants are some of the most popular algorithms employed in practice, especially the **no-U-turn sampler** (NUTS) (Hoffman and Gelman, 2014) which adaptively sets the integration time. In terms of complexity analysis, the paper Chen and Vempala (2019) (whose proof we followed in Theorem 5.2.1) provided a sharp analysis of ideal HMC. Other complexity results obtained for HMC under various assumptions include Mangoubi and Vishnoi (2018); Mangoubi and Smith (2019); Bou-Rabee et al. (2020, 2026). Discretization of Itô SDEs with multiplicative noise, as in Exercise 5.3, was considered in Erdogdu et al. (2018), in which such diffusions were used to establish some global non-convex optimization guarantees. An improved rate in KL divergence, under higher-order smoothness, was obtained in Li et al. (2025a).

The underdamped Langevin diffusion has been studied quantitatively in Cheng et al. (2018); Eberle et al. (2019); Dalalyan and Riou-Durand (2020); Ma et al. (2021); Zhang et al. (2023); Kim et al. (2025); Lehec (2025).

Besides randomized midpoint, the (idealized) **shifted ODE algorithm** (Foster et al., 2021) also achieves iteration complexity of $\tilde{O}(\kappa d^{1/3}/\varepsilon^{2/3})$ when applied to the underdamped Langevin diffusion. However, implementing the ODE requires the use of a numerical integrator, and it is currently not known how to preserve the rate when fully discretized. Moreover, the rate $\tilde{O}(\kappa d^{1/3}/\varepsilon^{2/3})$ is an optimal rate of discretization of the underdamped Langevin diffusion on the level of trajectories (Cao et al., 2021); see the Bibliographical Notes in Chapter 9.

Recent works have sharpened the results of Theorem 5.1.1 and Theorem 5.3.12. First, Kandasamy and Nagaraj (2024) obtained KL divergence guarantees (which are stronger than W_2 guarantees) in the LSI setting, with rates $\tilde{O}(\kappa^{3/2}d^{3/4}/\varepsilon)$ for RM-LMC and $\tilde{O}(\kappa^{17/12}d^{5/12}/\varepsilon^{1/2})$ for RM-ULMC. Then, Altschuler and Chewi (2026b) developed a KL version of the local error framework, leading to a KL divergence rate of $\tilde{O}(\kappa d^{1/2}/\varepsilon)$ for RM-LMC in the strongly log-concave setting, as well as results for the LSI and weakly log-concave ($\nabla^2 V \geq 0$) settings. The KL local error framework was extended to the underdamped Langevin diffusion in Altschuler et al. (2025), leading to the KL divergence rate $\tilde{O}(\kappa d^{1/3}/\varepsilon^{2/3})$ for RM-ULMC in the strongly log-concave setting, again with further results in the LSI and weakly log-concave settings. Combined with the accelerated hypocoercivity result (Theorem 5.3.16), it yields an iteration complexity of $\tilde{O}(\kappa^{5/6}d^{1/3}/\varepsilon^{2/3})$; see Theorem 5.3.17. However, at the time of publication, the entropic version of Theorem 5.3.16 was not available, so the original result incurred an additional dependence on $\log \chi^2(\mu_0 \parallel \pi)$. Finally, building on Kandasamy and Nagaraj (2024), Srinivasan and Nagaraj (2025) established the W_2 rates $\tilde{O}((\kappa^{4/3}d^{1/3} + \kappa d^{2/3})/\varepsilon^{2/3})$ for RM-LMC, and roughly $\tilde{O}(\kappa^{7/6}d^{1/3}/\varepsilon^{1/3})$ for RM-ULMC in the strongly log-concave setting.

The work of Zhang (2025) uses an anticipating Girsanov theorem and obtains KL and Rényi divergence bounds for non-adapted interpolations of discretizations, such as midpoint methods. Although this technique is not able to take advantage of the smaller weak error of randomized midpoint, it is well-suited for checking the cross-regularity assumptions of Altschuler and Chewi (2026b); Altschuler et al. (2025).

We remark that these rates are somewhat incomparable, and at the time of writing, the dust has not yet settled on the determination of the “correct” rates. The extension to KL divergence is also rather non-trivial; note, for instance, that a naïve adaptation of the methods of Sections 4.2 and 4.4 fails because RM-LMC does not admit a natural adapted interpolation.

The literature on hypocoercivity is large and we do not intend to comprehensively survey it here, but we give a few pointers. Our treatment of hypocoercivity is inspired by the monograph Villani (2009a), which gave a systematic general framework. The illuminating work of Baudoin (2017) relates the hypocoercive calculations to the Bakry–Émery calculus. The approach to hypocoercivity taken here is sometimes called “ H^1 hypocoercivity” because the constructed norm is equivalent to the $H^1(\pi)$ Sobolev norm, but there is another approach based on “ L^2 hypocoercivity”. The latter is called the “DMS framework” after Dolbeault et al. (2009, 2015); see Roussel and Stoltz (2018) for the application to underdamped Langevin. The DMS method was also adapted to piecewise deterministic Markov processes (PDMPs), such as the bouncy particle sampler, randomized Hamiltonian Monte Carlo, and the zigzag sampler, in Andrieu et al. (2021a), with further applications and extensions in Andrieu et al. (2021b); Dobson and Bierkens (2023). In the strongly log-concave case, another approach to show convergence in Rényi divergences is to establish a Harnack inequality for the underdamped Langevin diffusion, as in Exercise 3.4; see Guillin and Wang (2012); Altschuler et al. (2025).

In the case of Gaussian targets (and small perturbations thereof), the accelerated convergence rate was established in W_2 via a coupling argument in Schuh (2024). Then, building on the earlier work of Albritton et al. (2024), the work of Cao et al. (2023) established a sharp space-time Poincaré inequality. This was the first accelerated hypocoercivity result in the Poincaré setting. Subsequently, the proof was cast into the second-order lifting framework of Eberle and Lörler (2026), and we have followed their exposition in Section 5.3.6. Further developments include Brigati (2023); Eberle and Lörler (2024); Brigati et al. (2025); Eberle et al. (2025); Lehec (2025); Lu (2025). The space-time approach is also applicable to PDMPs (Lu and Wang, 2022b).

Next, it was observed in Fan et al. (2026) that the L^2 DMS method can also achieve the sharp L^2

result by modifying the form of the corrector (specifically, by introducing α into the definition of the operator $\mathcal{A}$). Building on this, J. Lu's breakthrough (Lu, 2026) established the sharp entropic decay; see Section 5.3.5. The proof was generalized to non-linear equations in Monmarché (2026), and to PDMPs such as randomized HMC in Monmarché and Wang (2026). The space-time LSI and accelerated hypercontractivity were developed in Li and Lu (2026). Most recently, the work of Mitra et al. (2026) gave an alternative approach to acceleration of randomized HMC (Theorem 5.2.6), which currently obtains the only known speed-up under log-concavity alone.

Exercises

Randomized Midpoint Discretization

⊢ Exercise 5.1 (Stochastic gradient LMC)

Assume, as usual, that $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. However, suppose that when we run LMC, we only have access to stochastic gradient evaluations. Namely, for each $x \in \mathbb{R}^d$, we can obtain a stochastic gradient $\hat{\nabla}V(x)$ which is unbiased, $\mathbb{E} \hat{\nabla}V(x) = \nabla V(x)$, and has uniformly bounded variance:

$$\mathbb{E}[\|\hat{\nabla}V(x) - \nabla V(x)\|^2] \leq \sigma^2 \quad \text{for all } x \in \mathbb{R}^d.$$

Use the local error framework to analyze this variant of LMC.

⊢ Exercise 5.2 (LMC under higher-order smoothness)

Assume that the potential satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and furthermore assume that $\|\nabla \Delta V\| \leq \zeta_0 + \zeta_1 \|\nabla V\|$. Define the scale-invariant quantities $\kappa := \beta/\alpha$, $\bar{\kappa}_0 := \zeta_0/\alpha^{3/2}$, and $\bar{\kappa}_1 := \zeta_1/\alpha$. Use the local error framework to show that LMC achieves iteration complexity $\tilde{O}((\bar{\kappa}_0 + (\kappa + \bar{\kappa}_1)\sqrt{\kappa d})/\varepsilon)$ in this case. In particular, if $\kappa, \bar{\kappa}_0, \bar{\kappa}_1 = O(1)$, then the iteration complexity becomes $\tilde{O}(d^{1/2}/\varepsilon)$, which should match what was obtained for the Gaussian case in Exercise 4.1. Note that if we have a Lipschitz condition in the operator norm, $\|\nabla^2 V(x) - \nabla^2 V(y)\|_{\text{op}} \leq \rho \|x - y\|$, then $\|\nabla \Delta V\| \leq \rho d^{1/2}$, so in this case it implies the iteration complexity $\tilde{O}((\kappa^{3/2} d^{1/2} + \rho d/\alpha^{3/2})/\varepsilon)$.

Hence, if the Hessian is Lipschitz in the operator norm, the complexity of LMC improves from $\tilde{O}(d/\varepsilon^2)$ to $\tilde{O}(d/\varepsilon)$, and under a more stringent derivative condition (which is satisfied by Gaussians), it further improves to $\tilde{O}(d^{1/2}/\varepsilon)$.

Hint: To bound the weak error, use Itô's formula.

⊢ Exercise 5.3 (Itô SDEs with multiplicative noise)

Consider a general Itô SDE of the form

$$dX_t = b(X_t) dt + \sigma(X_t) dB_t. \quad (5.E.1)$$

We consider the Euler–Maruyama discretization

$$\hat{X}_{(n+1)h} = \hat{X}_{nh} + h b(\hat{X}_{nh}) + \sigma(\hat{X}_{nh}) (B_{(n+1)h} - B_{nh}).$$

We adopt the following assumptions:

- 1 (Contractivity) $2 \langle b(x) - b(y), x - y \rangle + \|\sigma(x) - \sigma(y)\|_{\text{HS}}^2 \leq -2\alpha \|x - y\|^2$ for all $x, y \in \mathbb{R}^d$.
- 2 (Lipschitz coefficients) b is L -Lipschitz and σ is $\sqrt{L}$ -Lipschitz (in the $\|\cdot\|_{\text{HS}}$ norm).
- 3 (Coefficient bounds) $\|b(0)\| \leq B_b$ and $\|\sigma(0)\|_{\text{HS}} \leq B_\sigma$.

Writing π for the stationary distribution for (5.E.1) and $\kappa := L/\alpha$, use the local error framework to show that for all sufficiently small $\varepsilon > 0$ and for an appropriate choice of step size, we can obtain $\sqrt{\alpha} W_2(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$ in

$$N = O\left(\frac{\kappa^4}{\varepsilon^2} \left(\frac{B_b^2}{L} + B_\sigma^2\right) \log \frac{\alpha W_2^2(\mu_0, \pi)}{\varepsilon^2}\right) \quad \text{iterations}.$$

(This exercise is rather involved and also requires establishing second moment bounds in the spirit of Exercise 4.2.)

Hamiltonian Monte Carlo

Exercise 5.4 (Gaussian calculations for HMC)

Suppose that the target π is a standard Gaussian distribution. Compute the flow map F_t for the Hamiltonian dynamics. Also, if we start at the initial distribution $\mu_0 = \text{normal}(0, \sigma^2 I_d)$, show that the distribution μ_t over phase space at time t of ideal HMC is a Gaussian distribution, $\text{normal}(0, \Sigma_t)$, and compute $\Sigma_t \in \mathbb{R}^{2d \times 2d}$.

Exercise 5.5 (Basic properties of Hamiltonian dynamics)

In this exercise, we explore some fundamental properties of the Hamiltonian dynamics.

- 1 (Conservation of energy) Along the Hamiltonian dynamics, $(x_t, p_t)_{t \geq 0}$, show that $H(x_t, p_t) = H(x_0, p_0)$. In fact, for any function $f : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$ whose **Poisson bracket** with H vanishes, i.e.,

$$\{f, H\} := \langle \nabla_x f, \nabla_p H \rangle - \langle \nabla_p f, \nabla_x H \rangle = 0,$$

it holds that $f(x_t, p_t) = f(x_0, p_0)$.

- 2 (Conservation of volume) By differentiating $t \mapsto \det \nabla F_t(x, p)$ and using the flow map equation $\partial_t F_t(x, p) = \mathbf{J} \nabla H(F_t(x, p))$, prove that $\det \nabla F_t(x, p) = 1$ for all $t \geq 0$ and $x, p \in \mathbb{R}^d$. This shows that $F_t : \mathbb{R}^{2d} \rightarrow \mathbb{R}^{2d}$ is a volume-preserving map.
- 3 (Time reversibility) Suppose that $(x_t, p_t)_{t \in [0, T]}$ solve Hamilton's equations. Then, show that $(x_{T-t}, -p_{T-t})_{t \in [0, T]}$ also solve Hamilton's equations. In other words, if R is the momentum reversal operator, i.e.,

$$R = \begin{bmatrix} I_d & 0 \\ 0 & -I_d \end{bmatrix},$$

then $F_T^{-1} = R \circ F_T \circ R$.

Exercise 5.6 (Coercivity)

Prove Lemma 5.2.3.

Hint: Let $z := y - \frac{1}{\beta} \{\nabla f(y) - \nabla f(x)\}$. Apply the convexity inequality to $f(x) - f(z)$, and the smoothness inequality to $f(z) - f(y)$, in order to upper bound $f(x) - f(y)$. Combine this with the symmetric inequality for $f(y) - f(x)$.

The Underdamped Langevin Diffusion

Exercise 5.7 (Nesterov's algorithm in continuous time)

Consider the continuous-time formulation of Nesterov's algorithm, as given in Optimization Box 5.3.2. Assume that V is α -strongly convex and $p_0 = 0$. Prove the rate (5.3.3).

Hint: Let $z_t := x_t + \frac{2}{\gamma} p_t$. Consider the Lyapunov functional

$$\mathcal{L}_t := V(x_t) - V(x_\star) + \frac{\alpha}{2} \|z_t - x_\star\|^2.$$

Prove that $\dot{\mathcal{L}}_t \leq -\sqrt{\alpha} \mathcal{L}_t$. You may find the following identity to be helpful: $\langle z - x, z - x_\star \rangle = \frac{1}{2} (\|z - x\|^2 + \|z - x_\star\|^2 - \|x - x_\star\|^2)$.

▷ **Exercise 5.8** (Derivation of the ULMC updates)

Solve the SDE for ULMC to prove Lemma 5.3.9.

▷ **Exercise 5.9** (Movement bound for the underdamped Langevin diffusion)

Prove the movement bound for underdamped Langevin (Lemma 5.3.11).

▷ **Exercise 5.10** (Uniform-in-time LSI for ULMC)

Establish the analogue of Exercise 4.4 for ULMC. You may restrict attention to the contractive case $\alpha > 0$, and you can use the contraction for ULMC in the third part of Exercise 6.9.

▷ **Exercise 5.11** (Implementation of RM-ULMC)

Show how to implement RM-ULMC by computing $\text{cov}(\xi_{nh,s}^{(1)}, \xi_{nh,t}^{(1)})$, $\text{cov}(\xi_{nh,t}^{(1)}, \xi_{nh,h}^{(2)})$, and $\text{cov}(\xi_{nh,h}^{(2)})$.

▷ **Exercise 5.12** (Analysis of RM-ULMC)

In this exercise, we establish Theorem 5.3.12.

- 1 Prove an analogue of Lemma 5.1.6 for the underdamped Langevin diffusion (work in the twisted norm $\|\cdot\|$ throughout).
- 2 Show that if we take the expectation over u_n and v_n , denoted $\mathbb{E}_{u_n}$ and $\mathbb{E}_{v_n}$ respectively, then

$$\begin{aligned} \mathbb{E}_{u_n} \hat{X}_{(n+1)h} &= \hat{X}_{nh} + \frac{1 - \exp(-\gamma h)}{\gamma} \hat{P}_{nh} - \int_{nh}^{(n+1)h} \frac{1 - \exp(-\gamma((n+1)h - t))}{\gamma} \nabla V(\hat{X}_t^+) dt \\ &\quad + \xi_{nh,h}^{(1)}, \\ \mathbb{E}_{v_n} \hat{P}_{(n+1)h} &= \exp(-\gamma h) \hat{P}_{nh} - \int_{nh}^{(n+1)h} \exp(-\gamma((n+1)h - t)) \nabla V(\hat{X}_t^+) dt + \xi_{nh,h}^{(2)}. \end{aligned}$$

- 3 By comparing the expressions in the preceding part with the exact solution of the underdamped Langevin diffusion, establish the following local error bounds, assuming that $-\beta I_d \leq \nabla^2 V \leq \beta I_d$, $h \ll 1/\sqrt{\beta} \wedge \gamma/\beta$, and $\gamma \lesssim \sqrt{\beta}$:

$$\begin{aligned} h^{-1} \|\mathbb{E} \hat{X}_h - \mathbb{E} X_h\| + \|\mathbb{E} \hat{P}_h - \mathbb{E} P_h\| &\lesssim \beta^2 h^4 \|p\| + \beta^2 \gamma^{1/2} d^{1/2} h^{9/2} + \beta^2 h^5 \|\nabla V(x)\|, \\ h^{-1} \|\hat{X}_h - X_h\|_{L^2} + \|\hat{P}_h - P_h\|_{L^2} &\lesssim \beta h^2 \|p\| + \beta \gamma^{1/2} d^{1/2} h^{5/2} + \beta h^3 \|\nabla V(x)\|. \end{aligned}$$

Notice that the weak error is incredibly small!

- 4 Use the local error framework to prove Theorem 5.3.12.

▷ **Exercise 5.13** (Fokker–Planck equation for the underdamped Langevin diffusion)

Check the expression $\mathcal{L} = \gamma \mathcal{L}_{\text{OU}} + \mathcal{L}_{\text{Ham}}$ for the generator of the underdamped Langevin diffusion given in Section 5.3.4. Then, check the various calculations leading up to the Fokker–Planck equations (5.3.5) and (5.3.13).

▷ **Exercise 5.14** (Computations with adjoints and commutators)

Prove that $\mathcal{L}_{\text{Ham}}^* = -\mathcal{L}_{\text{Ham}}$ and that $\mathcal{L}_{\text{OU}} = -\nabla_p^* \nabla_p$, where the adjoints are taken in $L^2(\pi)$. Moreover, verify the commutator relations (Com).

▷ **Exercise 5.15** (Entropic hypocoercivity)

Adapt the proof of Theorem 5.3.14 to prove Theorem 5.3.15.

▷ **Exercise 5.16** (Linearizing the Lyapunov functions)

For $\mu_\varepsilon = (1 + \varepsilon \chi) \pi$, show that

$$\lim_{\varepsilon \searrow 0} \frac{\mathcal{L}_{\text{KL}}(\mu_\varepsilon)}{\varepsilon^2} = \frac{1}{2} \|\chi\|_{L^2(\pi)}^2 + c' \sqrt{\alpha} \langle \chi, (-\mathcal{L})^{-1} \Pi \mathcal{L}_{\text{Ham}} \chi \rangle_{L^2(\pi)}.$$

Thus, $\mathcal{L}_{\text{KL}}$ linearizes to a version of $\mathcal{L}_{\chi^2}$, where α is replaced with 0 in the definition of $\mathcal{A}$.

Hint: To linearize $T_{\mu_\varepsilon \rightarrow \pi}$, see Exercise 2.1. The following identity, which you should prove, could be helpful: $\langle p, \nabla_x \Pi f \rangle = \mathcal{L}_{\text{Ham}} \Pi f$.

▷ **Exercise 5.17** (Accelerated entropic hypocoercivity)

Here, we fill in some details for the proof of Theorem 5.3.16.

- 1 Prove (5.3.18).
- 2 Prove (5.3.19) as follows: let $n_{m, M^{-1}} := \text{normal}(m, M^{-1})$ and compute $\text{KL}(\mu \parallel n)$ by introducing the density $n_{m, M^{-1}}$ as an intermediary.

▷ **Exercise 5.18** (Relaxation time lower bound)

Here, we prove Theorem 5.3.22.

- 1 Using reversibility, show that $\text{relax}(\mathcal{L}) \asymp 1/\lambda$, where λ is the spectral gap of $\mathcal{L}$.
- 2 Let $f \in L^2(\pi)$ be such that $\int f d\pi = 0$, $\int f^2 d\pi = 1$, and $\mathcal{E}(f, f) = \lambda$. (Here, we assume the existence of an eigenfunction for simplicity, although the proof can avoid this.) Use the second-order lift property to show that $\|\mathcal{L}(f \circ \Pi)\|_{L^2(\pi)} \lesssim \sqrt{\lambda} \|f \circ \Pi\|_{L^2(\pi)}$.
- 3 Suppose that $(P_t)_{t \geq 0}$ satisfies the decay rate $\|P_t f\|_{L^2(\pi)} \leq \exp(-\omega(t - t_0)) \|f\|_{L^2(\pi)}$ for all $t \geq 0$ and all $f \in L^2(\pi)$ with $\int f d\pi = 0$. Using the fact that $(-\mathcal{L})^{-1} = \int_0^\infty P_t dt$ (justify!) and by considering $\int_0^\infty \|P_t f\|_{L^2(\pi)} dt$, prove that $t_0 + \omega^{-1} \gtrsim 1/\sqrt{\lambda}$.
- 4 Finally, show that the decay in the preceding part holds with $\omega^{-1} \asymp t_0 = \text{relax}(\mathcal{L})$. Combine everything together to deduce the result.

▷ **Exercise 5.19** (Decay under a space-time Poincaré inequality)

Prove that (5.3.25) implies (5.3.26).

▷ **Exercise 5.20** (Space-time Poincaré inequality without log-concavity)

In the first part of Lemma 5.3.23, how is the bound affected if we do not assume log-concavity, but instead assume that $\nabla^2 V \geq -\beta I_d$ for some $\beta > 0$? How does this affect the final convergence estimate?

▷ **Exercise 5.21** (Randomized Hamiltonian Monte Carlo)

Consider ideal Hamiltonian Monte Carlo with randomized integration times, namely, the integration times are drawn i.i.d. with the exponential distribution with rate γ .

- 1 Argue that the generator is given by $\gamma(\Pi - \text{id}) + \mathcal{L}_{\text{Ham}}$.

- 2 Adapt the proof of the space-time Poincaré inequality to randomized Hamiltonian Monte Carlo. Namely, explain how to substitute for the inequalities leading up to (5.3.27).

▷ **Exercise 5.22 (Space-time LSI)**

In this exercise, we prove Theorem 5.3.28.

- 1 Rewrite (5.3.29) as $(\partial_t + \mathcal{L}_{\text{Ham}}) \rho_t = \gamma \mathcal{L}_{\text{OU}} \rho_t + \nabla_p^*(\rho_t v_t)$, where $v_t := u_t + \gamma \nabla_p \log \rho_t$. Let $\bar{v}_t := \mathbb{E}_{\mu_t}[v_t(X_t, P_t) \mid X_t = \cdot]$. Show that

$$\begin{aligned} \partial_t \text{KL}(\mu_t \parallel \pi) &= -\gamma \text{Fl}^P(\mu_t \parallel \pi) + \int \langle v_t, \nabla_p \log \rho_t \rangle d\mu_t, \\ \partial_t \mathcal{C}(\mu_t) &\leq \text{Fl}^P(\mu_t \parallel \pi) - \gamma \mathcal{C}(\mu_t) - \text{KL}(\mu_t^X \parallel \pi) + \int \langle \bar{v}_t, \text{id} - T_{\mu_t^X} \rangle d\mu_t^X. \end{aligned}$$

- 2 Deduce that if we choose $\gamma = \sqrt{\alpha}$, then

$$\partial_t \mathcal{L}_{\text{KL}}(\mu_t) \leq -\frac{c' \sqrt{\alpha}}{2} \text{KL}(\mu_t \parallel \pi) + O\left(\frac{1}{\sqrt{\alpha}}\right) \mathbb{E}_{\mu_t}[\|v_t\|^2].$$

Integrate this to obtain

$$\sqrt{\alpha} \int_0^T \text{KL}(\mu_t \parallel \pi) dt \lesssim \text{KL}(\mu_0 \parallel \pi) + \sqrt{\alpha} \int_0^T \text{Fl}^P(\mu_t \parallel \pi) dt + \frac{1}{\sqrt{\alpha}} \int_0^T \mathbb{E}_{\mu_t}[\|u_t\|^2] dt.$$

- 3 Note that $\text{KL}(\mu_0 \parallel \pi) \leq \text{KL}(\mu_t \parallel \pi) + \int_0^t |\partial_s \text{KL}(\mu_s \parallel \pi)| ds$. Average this inequality over $t \in [0, T]$, substitute this into the inequality above, and establish Theorem 5.3.28.

▷ **Exercise 5.23 (Accelerated hypercontractivity)**

Let $(P_t)_{t \geq 0}$, $(P_t^*)_{t \geq 0}$ denote the underdamped Langevin semigroup and its $L^2(\pi)$ dual, respectively. The goal of this exercise is to establish the hypercontractivity inequality $\|P_T^* f\|_{L^{q_1}(\pi)} \leq \|f\|_{L^{q_0}(\pi)}$ for all f and some $q_1 > q_0$. Due to the hypocoercive nature of the semigroup, the simple approach of Exercise 2.10 does not apply.

- 1 By duality, show that this is equivalent to the following statement: for all pairs of positive and smooth functions (φ, ψ) , normalized so that $\|\varphi\|_{L^{q_0}(\pi)} = \|\psi\|_{L^{q_1^*}(\pi)} = 1$, it holds that $Z := \int P_T^* \varphi \psi d\pi \leq 1$. Here, for an exponent $q \geq 1$, we denote by $q^* := q/(q-1)$ its Hölder conjugate.
- 2 For $0 \leq t \leq T$, let $\varphi_t := P_t^* \varphi$, $\psi_t := P_{T-t} \psi$, and note that $Z = \int \varphi_t \psi_t d\pi$. Thus, for each $t \in [0, T]$, we can define a measure μ_t such that $\rho_t := \frac{d\mu_t}{d\pi} = \varphi_t \psi_t / Z$. Show that $(\rho_t)_{t \in [0, T]}$ satisfies (5.3.29) with $u_t = \gamma (\nabla_p \log \psi_t - \nabla_p \log \varphi_t)$.
- 3 Prove the key identity

$$2 \log Z = 2 \int \rho_0 \log \varphi d\pi + 2 \int \rho_T \log \psi d\pi - \text{KL}(\mu_0 \parallel \pi) - \text{KL}(\mu_T \parallel \pi) - \gamma \mathcal{D}_T(\varphi, \psi),$$

where

$$\mathcal{D}_T(\varphi, \psi) := \int_0^T \mathbb{E}_{\mu_t}[\|\nabla_p \log \varphi_t\|^2 + \|\nabla_p \log \psi_t\|^2] dt.$$

- 4 Apply the space-time log-Sobolev inequality (Theorem 5.3.28) to show that, as long as $T \gg 1/\sqrt{\alpha}$, we have $\int_0^T \text{KL}(\mu_t \parallel \pi) dt \lesssim \mathcal{D}_T(\varphi, \psi)$. Further deduce that

$$\text{KL}(\mu_0 \parallel \pi) + \text{KL}(\mu_T \parallel \pi) \lesssim (\gamma + T^{-1}) \mathcal{D}_T(\varphi, \psi).$$

- 5 Show that $\int \rho_0 \log \varphi \, d\pi \leq q_0^{-1} \text{KL}(\mu_0 \parallel \pi)$ and $\int \rho_T \log \psi \, d\pi \leq (q_1^*)^{-1} \text{KL}(\mu_T \parallel \pi)$. Putting everything together, show that by choosing $T \asymp 1/\sqrt{\alpha}$, there are exponents $\bar{q}_0 < 2 < \bar{q}_1$ such that $\|\mathbf{P}_T^* f\|_{L^{\bar{q}_1}(\pi)} \leq \|f\|_{L^{\bar{q}_0}(\pi)}$.
- 6 The above result establishes hypercontractivity for one pair of exponents $(\bar{q}_0, \bar{q}_1)$ at the time scale $T \asymp 1/\sqrt{\alpha}$. By Riesz–Thorin interpolation, it implies that for all $q_1 > q_0$ with $q_1 - 1 \leq \bar{C}(q_0 - 1)$, it holds that $\|\mathbf{P}_T^* f\|_{L^{q_1}(\pi)} \leq \|f\|_{L^{q_0}(\pi)}$. This can also be iterated and averaged over time, yielding the full statement of [Theorem 5.3.30](#). As an optional exercise, prove these assertions.

Convergence in Rényi Divergence

In this chapter, we study sampling guarantees which hold in Rényi divergences. Recall from Section 2.2.5 that the Rényi divergences are a family of information divergences, indexed by a parameter q , such that the Rényi divergence of order $q = 1$ is the KL divergence, and the Rényi divergence of order 2 is related to the chi-squared divergence via $\mathcal{R}_2 = \log(1 + \chi^2)$. Rényi divergence guarantees are stronger than W_2 or KL guarantees, and they have been of interest in their own right in differential privacy (Mironov, 2017). The results from this chapter will also be used in Chapter 7.

In Section 2.2.5, we studied the continuous-time convergence of the Langevin diffusion in Rényi divergence under either a Poincaré inequality or a log-Sobolev inequality. We will build upon these results in order to study the discretized LMC algorithm. Then, we will consider the underdamped Langevin diffusion, which was introduced in Section 5.3.

6.1 Analysis of LMC via Interpolation Argument

In this section, we follow Chewi et al. (2025a), which generalizes the interpolation argument of Theorem 4.2.7 and provides a clean Rényi convergence proof for LMC under the assumption of a log-Sobolev inequality (LSI).

As in Section 4.2, we begin by writing a differential inequality for the Rényi divergence along the interpolation (4.2.3) of LMC. The proof is a combination of the proofs of Theorem 2.2.25 and Corollary 4.2.5, so it is left as Exercise 6.1. Throughout this section, let $q \geq 2$ be fixed.

Proposition 6.1.1. *Along the law $(\hat{\mu}_t)_{t \geq 0}$ of the interpolated process (4.2.3),*

$$\partial_t \mathcal{R}_q(\hat{\mu}_t \parallel \pi) \leq -\frac{3}{q} \frac{\mathbb{E}_\pi[\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)} + q \mathbb{E}[\psi_t(\hat{X}_t) \|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_\infty)\|^2],$$

where $\rho_t := \frac{d\hat{\mu}_t}{d\pi}$ and $\psi_t := \rho_t^{q-1} / \mathbb{E}_\pi(\rho_t^q)$.

In analogy with the usual Fisher information, the quantity

$$\text{FI}_q(\mu \parallel \pi) := \frac{4}{q} \frac{\mathbb{E}_\pi[\|\nabla(\rho^{q/2})\|^2]}{\mathbb{E}_\pi(\rho^q)}, \quad \rho := \frac{d\mu}{d\pi},$$

may be considered the “Rényi Fisher information”. As in the proof of [Theorem 2.2.25](#), under a log-Sobolev inequality, we have

$$\text{FI}_q(\mu \parallel \pi) \geq \frac{2}{qC_{\text{LSI}}} \mathcal{R}_q(\mu \parallel \pi),$$

so the first term in [Proposition 6.1.1](#) provides a decay in the Rényi divergence. In the discretization analysis, our task is to control the second term.

Note that

$$\mathbb{E} \psi_t(\hat{X}_t) = \mathbb{E}_{\hat{\mu}_t} \psi_t = \mathbb{E}_\pi \left[\rho_t \frac{\rho_t^{q-1}}{\mathbb{E}_\pi(\rho_t^q)} \right] = 1,$$

so $\psi_t(\hat{X}_t)$ acts as a change of measure. The main difficulty of the proof is that whereas we know how to control the term $\|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_{t-})\|^2$ under the original probability measure $\mathbb{P}$ (indeed, this is precisely what we accomplished in [Theorem 4.2.7](#)), it is not straightforward to control this term under the measure $\mathbf{P}$ defined by $\frac{d\mathbf{P}}{d\mathbb{P}} = \psi_t(\hat{X}_t)$. Towards this end, we shall employ change of measure inequalities that allow us to relate expectations under $\mathbb{P}$ to expectations under $\mathbf{P}$. Note that $\psi_t = 1$ when $q = 1$, which is why these difficulties can be avoided when working with the KL divergence.

The main theorem that we wish to prove is as follows.

Theorem 6.1.2 ([Chewi et al. \(2025a\)](#)). *For $n \in \mathbb{N}$, let $\hat{\mu}_{nh}$ denote the law of the n -th iterate of LMC with step size $h > 0$. Assume that the target $\pi \propto \exp(-V)$ satisfies $C_{\text{LSI}}(\pi) \leq 1/\alpha$ and that ∇V is β -Lipschitz. Then, for all $h \leq 1/192\beta\kappa q^2$ and $N \geq N_0$, it holds that*

$$\mathcal{R}_q(\hat{\mu}_{Nh} \parallel \pi) \leq \exp\left(-\frac{\alpha(N - N_0)h}{4}\right) \mathcal{R}_2(\hat{\mu}_0 \parallel \pi) + \tilde{O}\left(\frac{\beta^2 dhq}{\alpha}\right),$$

where $N_0 := \lceil \frac{2}{\alpha h} \log(q - 1) \rceil$. In particular, for all $\varepsilon \in [0, \sqrt{d/q}]$, if we choose the step size $h = \tilde{\Theta}(\varepsilon^2/\beta\kappa dq)$, then we obtain the guarantee $\sqrt{\mathcal{R}_q(\hat{\mu}_{Nh} \parallel \pi)} \leq \varepsilon$ after

$$N = \tilde{O}\left(\frac{\kappa^2 dq}{\varepsilon^2} \log \mathcal{R}_2(\hat{\mu}_0 \parallel \pi)\right) \quad \text{iterations}.$$

For clarity of exposition, we begin with a discretization analysis that incurs a worse dependence on q . Afterwards, we show how to improve the dependence on q via a hypercontractivity argument.

Proof of Theorem 6.1.2 with suboptimal dependence on q As in the proof of [Theorem 4.2.7](#), our aim is to control the error term $\mathbb{E}[\psi_t(\hat{X}_t) \|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_{t-})\|^2]$, where from the β -smoothness of V and from $h \leq 1/3\beta$ we have

$$\|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_{t-})\|^2 \leq 9\beta^2 (t - t_-)^2 \|\nabla V(\hat{X}_t)\|^2 + 6\beta^2 \|B_t - B_{t-}\|^2.$$

There are two terms to control. For the first term, applying the duality lemma for the Fisher

information (Lemma 4.2.6) to the measure $\psi_t \hat{\mu}_t$,

$$\begin{aligned} \mathbb{E}_{\psi_t \hat{\mu}_t} [\|\nabla V\|^2] &\leq \text{FI}(\psi_t \hat{\mu}_t \parallel \pi) + 2\beta d = \mathbb{E}_{\hat{\mu}_t} \left[\psi_t \left\| \nabla \log \left(\psi_t \frac{d\hat{\mu}_t}{d\pi} \right) \right\|^2 \right] + 2\beta d \\ &= \frac{\mathbb{E}_\pi [\rho_t^q \|\nabla \log(\rho_t^q)\|^2]}{\mathbb{E}_\pi(\rho_t^q)} + 2\beta d = \frac{4 \mathbb{E}_\pi [\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)} + 2\beta d, \end{aligned}$$

where we used the identity

$$\mathbb{E}_{\hat{\mu}_t} \left[\psi_t \left\| \nabla \log \left(\psi_t \frac{d\hat{\mu}_t}{d\pi} \right) \right\|^2 \right] = \frac{4 \mathbb{E}_\pi [\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)} \quad (6.1.3)$$

which follows from the ordinary chain rule.

For the second error term, we must control the term $\|B_t - B_{t_-}\|^2$ under the measure $\mathbf{P}$, where $\frac{d\mathbf{P}}{d\mathbb{P}} = \psi_t(\hat{X}_t)$. The difficulty is that under $\mathbf{P}$, B is no longer a standard Brownian motion, so it is difficult to control this term directly. Instead, we apply the Donsker–Varadhan variational principle (Theorem 1.5.7) to relate the expectation under $\mathbf{P}$ (denoted $\mathbf{E}$) with the expectation under $\mathbb{P}$. For any random variable ζ , it yields

$$\mathbf{E}\zeta \leq \text{KL}(\mathbf{P} \parallel \mathbb{P}) + \log \mathbb{E} \exp \zeta.$$

Applying this to $\zeta := c (\|B_t - B_{t_-}\| - \mathbb{E}\|B_t - B_{t_-}\|)^2$ for a constant $c > 0$ to be chosen later, we obtain

$$\begin{aligned} \mathbf{E}[\|B_t - B_{t_-}\|^2] &\leq 2 \mathbb{E}[\|B_t - B_{t_-}\|^2] + \frac{2}{c} \mathbf{E}\zeta \\ &\leq 2d(t - t_-) + \frac{2}{c} \left\{ \text{KL}(\mathbf{P} \parallel \mathbb{P}) + \log \mathbb{E} \exp(c (\|B_t - B_{t_-}\| - \mathbb{E}\|B_t - B_{t_-}\|)^2) \right\}. \end{aligned}$$

Under $\mathbb{P}$, $B_t - B_{t_-} \sim \text{normal}(0, (t - t_-) I_d)$. Applying concentration of measure for the Gaussian distribution (see, e.g., Theorem 2.4.3), if $c \lesssim 1/t_{-}$, then

$$\mathbb{E} \exp(c (\|B_t - B_{t_-}\| - \mathbb{E}\|B_t - B_{t_-}\|)^2) \leq 2.$$

In fact, it suffices to take $c = 1/8(t - t_-)$. Next, using the LSI for π ,

$$\begin{aligned} \text{KL}(\mathbf{P} \parallel \mathbb{P}) &= \mathbb{E}_{\psi_t \hat{\mu}_t} \log \psi_t = \mathbb{E}_{\psi_t \hat{\mu}_t} \log \frac{\rho_t^{q-1}}{\mathbb{E}_{\hat{\mu}_t}(\rho_t^{q-1})} = \frac{q-1}{q} \mathbb{E}_{\psi_t \hat{\mu}_t} \log \frac{\rho_t^q}{\mathbb{E}_{\hat{\mu}_t}(\rho_t^{q-1})^{q/(q-1)}} \\ &= \frac{q-1}{q} \left\{ \mathbb{E}_{\psi_t \hat{\mu}_t} \log \frac{\rho_t^q}{\mathbb{E}_{\hat{\mu}_t}(\rho_t^{q-1})} - \underbrace{\frac{1}{q-1} \log \mathbb{E}_{\hat{\mu}_t}(\rho_t^{q-1})}_{\geq 0} \right\} \\ &\leq \frac{q-1}{q} \text{KL}(\psi_t \hat{\mu}_t \parallel \pi) \\ &\leq \frac{q-1}{2\alpha q} \mathbb{E}_{\psi_t \hat{\mu}_t} \left[\left\| \nabla \log \left(\psi_t \frac{d\hat{\mu}_t}{d\pi} \right) \right\|^2 \right] = \frac{2(q-1)}{\alpha q} \frac{\mathbb{E}_\pi [\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)}, \end{aligned}$$

where we applied the identity (6.1.3). Hence,

$$\begin{aligned} & \mathbb{E}[\psi_t(\hat{X}_t) \|B_t - B_{t_-}\|^2] \\ & \leq 2d(t - t_-) + \frac{32h(q-1)}{\alpha q} \frac{\mathbb{E}_\pi[\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)} + (16 \log 2)(t - t_-) \\ & \leq 14d(t - t_-) + \frac{32h}{\alpha} \frac{\mathbb{E}_\pi[\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)}. \end{aligned}$$

All in all, applying Proposition 6.1.1 and collecting the error terms,

$$\begin{aligned} \partial_t \mathcal{R}_q(\hat{\mu}_t \| \pi) & \leq -\frac{3}{q} \frac{\mathbb{E}_\pi[\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)} + 9\beta^2 q (t - t_-)^2 \left\{ \frac{4\mathbb{E}_\pi[\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)} + 2\beta d \right\} \\ & \quad + 6\beta^2 q \left\{ 14d(t - t_-) + \frac{32h}{\alpha} \frac{\mathbb{E}_\pi[\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)} \right\}. \end{aligned}$$

From $h \leq 1/192\beta\kappa q^2$, we can absorb some of the error terms into the decay term and apply the LSI for π (see Theorem 2.2.25), yielding

$$\begin{aligned} \partial_t \mathcal{R}_q(\hat{\mu}_t \| \pi) & \leq -\frac{1}{q} \frac{\mathbb{E}_\pi[\|\nabla(\rho_t^{q/2})\|^2]}{\mathbb{E}_\pi(\rho_t^q)} + O(\beta^3 dh^2 q + \beta^2 dhq) \\ & \leq -\frac{\alpha}{2q} \mathcal{R}_q(\hat{\mu}_t \| \pi) + O(\beta^2 dhq). \end{aligned}$$

This implies the differential inequality

$$\partial_t \left\{ \exp\left(\frac{\alpha(t - t_-)}{2q}\right) \mathcal{R}_q(\hat{\mu}_t \| \pi) \right\} \lesssim \exp\left(\frac{\alpha(t - t_-)}{2q}\right) \beta^2 dhq \lesssim \beta^2 dhq.$$

Integrating this over $t \in [nh, (n+1)h]$ yields

$$\mathcal{R}_q(\hat{\mu}_{(n+1)h} \| \pi) \leq \exp\left(-\frac{\alpha h}{2q}\right) \mathcal{R}_q(\hat{\mu}_{nh} \| \pi) + O(\beta^2 dh^2 q).$$

Unrolling the recursion,

$$\mathcal{R}_q(\hat{\mu}_{Nh} \| \pi) \leq \exp\left(-\frac{\alpha Nh}{2q}\right) \mathcal{R}_q(\hat{\mu}_0 \| \pi) + O\left(\frac{\beta^2 dhq^2}{\alpha}\right). \quad \square$$

We pause to reflect upon the proof. As discussed above, the key steps are to use change of measure inequalities in order to relate expectations under $\mathbf{P}$ to expectations under $\mathbb{P}$. This is accomplished via the Fisher information duality lemma (Lemma 4.2.6) and the Donsker–Varadhan variational principle (Theorem 1.5.7). These inequalities yield an additional error term of the form $\text{FI}(\psi_t \hat{\mu}_t \| \pi)$ (for the latter, this error term appears after an application of the LSI). The magical part of the calculation is that $\text{FI}(\psi_t \hat{\mu}_t \| \pi)$ is precisely equal to the Rényi Fisher information (up to constants), and when the step size h is sufficiently small it can be absorbed into the decay term of the differential inequality in Proposition 6.1.1.

The proof above implies an iteration complexity whose dependence on q scales as $N = O(q^3)$. In order to improve the dependence on q , we modify the differential inequality of Proposition 6.1.1 by making the parameter q time-dependent, similarly to the hypercontractivity principle (Exercise 2.10). The proof is left as Exercise 6.1.

Proposition 6.1.4 (Hypercontractivity). *Suppose that $C_{\text{LSI}}(\pi) \leq 1/\alpha$. Along the law $(\hat{\mu}_t)_{t \geq 0}$ of the interpolated process (4.2.3), if we define the parameter $q(t) := 1 + (q_0 - 1) \exp \frac{\alpha t}{2}$, then*

$$\begin{aligned} \partial_t \left(\frac{1}{q(t)} \log \int \rho_t^{q(t)} d\pi \right) \leq & -\frac{2(q(t) - 1)}{q(t)^2} \frac{\mathbb{E}_\pi[\|\nabla(\rho_t^{q(t)/2})\|^2]}{\mathbb{E}_\pi(\rho_t^{q(t)})} \\ & + (q(t) - 1) \mathbb{E}[\psi_t(\hat{X}_t) \|\nabla V(\hat{X}_t) - \nabla V(\hat{X}_{t-})\|^2], \end{aligned}$$

where $\rho_t := \frac{d\hat{\mu}_t}{d\pi}$ and $\psi_t := \rho_t^{q(t)-1} / \mathbb{E}_\pi(\rho_t^{q(t)})$.

Proof of Theorem 6.1.2 with improved dependence on q Let $\bar{q} \geq 3$.

Initial waiting phase. We apply hypercontractivity (Proposition 6.1.4) with $q_0 = 2$ and for $t \leq N_0 h$, where $N_0 := \lceil \frac{2}{\alpha h} \log(\bar{q} - 1) \rceil$. Note that $\bar{q} \leq q(N_0 h) \leq 2\bar{q}$. The bound on the error term from the previous proof yields

$$\partial_t \left(\frac{1}{q(t)} \log \int \rho_t^{q(t)} d\pi \right) \lesssim \beta^2 dh q(t).$$

Integrating this over $t \in [nh, (n+1)h]$ yields

$$\frac{1}{q((n+1)h)} \log \int \rho_{(n+1)h}^{q((n+1)h)} d\pi - \frac{1}{q(nh)} \log \int \rho_{nh}^{q(nh)} d\pi \lesssim \beta^2 dh^2 \bar{q}.$$

Unrolling the recursion yields

$$\frac{1}{q(N_0 h)} \log \int \rho_{N_0 h}^{q(N_0 h)} d\pi - \frac{1}{2} \log \int \rho_0^2 d\pi \lesssim \beta^2 dh^2 \bar{q} N_0 \leq \tilde{O}\left(\frac{\beta^2 dh \bar{q}}{\alpha}\right).$$

Finishing the convergence analysis. Next, after shifting time indices and applying the previous proof of Theorem 6.1.2 with $q = 2$,

$$\begin{aligned} \mathcal{R}_{\bar{q}}(\hat{\mu}_{(N+N_0)h} \parallel \pi) & \leq \frac{1}{q(N_0 h) - 1} \log \int \rho_{(N+N_0)h}^{q(N_0 h)} d\pi \leq \frac{3}{4} \log \int \rho_{N_0 h}^2 d\pi + \tilde{O}\left(\frac{\beta^2 dh \bar{q}}{\alpha}\right) \\ & \leq \frac{3}{4} \exp\left(-\frac{\alpha N h}{4}\right) \mathcal{R}_2(\hat{\mu}_0 \parallel \pi) + \tilde{O}\left(\frac{\beta^2 dh \bar{q}}{\alpha}\right). \end{aligned}$$

This proves the desired result. $\square$

The proof of Theorem 6.1.2 is rather specific to the LSI case because we use the LSI to bound the KL term $\text{KL}(\psi_t \hat{\mu}_t \parallel \pi)$ via the Rényi Fisher information, which is then absorbed into the differential inequality of Proposition 6.1.1. However, it turns out that rather than assuming an LSI for π , it suffices to have an LSI for μ_t for all $t \geq 0$ (possibly with an LSI constant that grows with t). One situation in which this holds is when we initialize LMC with a measure μ_0 that satisfies an LSI, and the potential V is convex. Note that this situation is not included in the case when π satisfies an LSI, because V may only have linear growth at infinity (whereas from Theorem 2.4.3, if $\pi \propto \exp(-V)$ satisfies an LSI, then V necessarily has quadratic growth at infinity). We explore this in Exercise 6.2.

6.2 Analysis of LMC via Girsanov's Theorem

We now provide a Rényi analysis for LMC based on Girsanov's theorem; see Sections 3.2.2 and 4.4 for background. The advantage of this approach is that it is more broadly applicable, as it is less reliant

on seemingly miraculous calculations. The results of this section will also be used for the analysis of MALA in Chapter 7.

Discretization error.

Following the proof of Theorem 4.4.1 in Section 4.4, let $\mathbf{P}_T, \mathbf{W}_T$ be path measures such that under $\mathbf{P}_T$, $(X_t)_{t \in [0, T]}$ is the interpolated LMC process, and under $\mathbf{W}_T$, $(X_t)_{t \in [0, T]}$ is the Langevin diffusion. Both processes are initialized at some measure $\hat{\mu}_0$. Our goal is to prove the following theorem.

Theorem 6.2.1 (LMC Rényi discretization bound). *Assume that the potential V is β -smooth and that $h \lesssim 1/\beta^2 q^2 T$ where $T := Nh$. Then, for all $q \geq 2$,*

$$\mathcal{R}_q(\mathbf{P}_T \parallel \mathbf{W}_T) \lesssim \frac{1}{q} \max_{n=0,1,\dots,N-1} \log \mathbb{E}^{\mathbf{W}_T} \exp \left\{ O(\beta^2 h^2 q^2 T \|\nabla V(X_{nh})\|^2) \right\} + \beta^2 dhqT.$$

Proof From Girsanov's theorem (Theorem 3.2.8; see also Section 4.4), we have¹

$$\begin{aligned} \mathbb{E}^{\mathbf{W}_T} \left[\left(\frac{d\mathbf{P}_T}{d\mathbf{W}_T} \right)^q \right] &= \mathbb{E}^{\mathbf{W}_T} \exp \left(\frac{q}{\sqrt{2}} \int_0^T \langle \nabla V(X_{t-}) - \nabla V(X_t), dB_t \rangle \right. \\ &\quad \left. - \frac{q}{4} \int_0^T \|\nabla V(X_{t-}) - \nabla V(X_t)\|^2 dt \right). \end{aligned}$$

Unlike the analysis in KL divergence, the stochastic integral is now inside the exponential. The first step is to remove this term, which we accomplish as follows. First, by the Cauchy–Schwarz inequality and for any $\lambda > 0$,

$$\begin{aligned} \left\{ \mathbb{E}^{\mathbf{W}_T} \left[\left(\frac{d\mathbf{P}_T}{d\mathbf{W}_T} \right)^q \right] \right\}^2 &\leq \mathbb{E}^{\mathbf{W}_T} \exp \left(\sqrt{2} q \int_0^T \langle \nabla V(X_{t-}) - \nabla V(X_t), dB_t \rangle \right. \\ &\quad \left. - \lambda \int_0^T \|\nabla V(X_{t-}) - \nabla V(X_t)\|^2 dt \right) \\ &\quad \times \mathbb{E}^{\mathbf{W}_T} \exp \left(\left(\lambda - \frac{q}{2} \right) \int_0^T \|\nabla V(X_{t-}) - \nabla V(X_t)\|^2 dt \right). \end{aligned}$$

Next, recall from Section 3.2.2 that for any martingale M , the exponential martingale $\mathcal{E}(M) := \exp(M - \frac{1}{2} [M, M])$ is a local martingale. We apply this to the martingale M given by $M_t := \sqrt{2} q \int_0^t \langle \nabla V(X_{s-}) - \nabla V(X_s), dB_s \rangle$ by taking $\lambda = q^2$. Then, $\mathcal{E}(M)$ is a non-negative local martingale, so it is a supermartingale, and $\mathbb{E}^{\mathbf{W}_T} \mathcal{E}(M)_T \leq 1$. Finally, noting that the expectation above is at least 1, we can drop the square and we obtain

$$\begin{aligned} \mathbb{E}^{\mathbf{W}_T} \left[\left(\frac{d\mathbf{P}_T}{d\mathbf{W}_T} \right)^q \right] &\leq \mathbb{E}^{\mathbf{W}_T} \exp \left(q^2 \int_0^T \|\nabla V(X_{t-}) - \nabla V(X_t)\|^2 dt \right) \\ &\leq \mathbb{E}^{\mathbf{W}_T} \exp \left(\beta^2 q^2 \int_0^T \|X_t - X_{t-}\|^2 dt \right). \end{aligned}$$

Let us first consider a single iteration, over the time interval $[0, h]$, and suppose that the processes

¹ Again, we ignore Novikov's condition, which can be avoided via localization.

are started at $x \in \mathbb{R}^d$. Then,

$$\begin{aligned} \|X_t - x\| &= \left\| - \int_0^t \nabla V(X_s) ds + \sqrt{2} B_t \right\| \leq \int_0^t \|\nabla V(X_s)\| ds + \sqrt{2} \|B_t\| \\ &\leq \beta \int_0^t \|X_s - x\| ds + h \|\nabla V(x)\| + \sqrt{2} \|B_t\|. \end{aligned}$$

By Grönwall's inequality, if $h \leq 1/\beta$, it yields, for all $t \in [0, h]$,

$$\|X_t - x\| \leq 3h \|\nabla V(x)\| + 3\sqrt{2} \sup_{s \in [0, t]} \|B_s\|. \quad (6.2.2)$$

Therefore, for any $\eta > 0$,

$$\mathbb{E}^{\mathbf{W}_T} \exp\left(\eta \sup_{t \in [0, h]} \|X_t - x\|^2\right) \leq \exp(18\eta h^2 \|\nabla V(x)\|^2) \mathbb{E}^{\mathbf{W}_T} \exp\left(36\eta \sup_{t \in [0, h]} \|B_t\|^2\right).$$

We require a tail bound for Brownian motion ([Lemma 6.2.4](#) below), which yields

$$\log \mathbb{E}^{\mathbf{W}_T} \exp\left(\eta \sup_{t \in [0, h]} \|X_t - x\|^2\right) \lesssim \eta h^2 \|\nabla V(x)\|^2 + \eta dh, \quad (6.2.3)$$

provided that $\eta \lesssim 1/h$.

We must now iterate this bound. One approach is to condition on the process up until time $(N-1)h$ and apply (6.2.3), which yields

$$\begin{aligned} &\mathbb{E}^{\mathbf{W}_T} \exp\left(\beta^2 q^2 \int_0^{Nh} \|X_t - X_{t-}\|^2 dt\right) \\ &\leq \mathbb{E}^{\mathbf{W}_T} \exp\left(\beta^2 q^2 \int_0^{(N-1)h} \|X_t - X_{t-}\|^2 dt + O(\beta^2 h^3 q^2 \|\nabla V(X_{(N-1)h})\|^2 + \beta^2 dh^2 q^2)\right) \end{aligned}$$

but now the terms in the exponential are dependent and we cannot peel off the next term. To circumvent this issue, we instead apply Jensen's inequality (or equivalently, the AM–GM inequality):

$$\mathbb{E}^{\mathbf{W}_T} \exp\left(\beta^2 q^2 \int_0^T \|X_t - X_{t-}\|^2 dt\right) \leq \frac{1}{T} \int_0^T \mathbb{E}^{\mathbf{W}_T} \exp(\beta^2 q^2 T \|X_t - X_{t-}\|^2) dt.$$

Applying (6.2.3) conditionally, it yields

$$\mathbb{E}^{\mathbf{W}_T} \left[\left(\frac{d\mathbf{P}_T}{d\mathbf{W}_T} \right)^q \right] \leq \frac{1}{T} \int_0^T \mathbb{E}^{\mathbf{W}_T} \exp\left(O(\beta^2 h^2 q^2 T \|\nabla V(X_{t-})\|^2 + \beta^2 dhq^2 T)\right) dt$$

provided that $\beta^2 q^2 T \lesssim 1/h$, i.e., $h \lesssim 1/\beta^2 q^2 T$. The result follows. $\square$

In the proof, we used the following lemma; see [Exercise 6.4](#).

Lemma 6.2.4 (Brownian tail bound). *Let $(B_t)_{t \geq 0}$ denote a standard Brownian motion in $\mathbb{R}^d$. Then, for $\lambda, t > 0$ such that $\lambda < \frac{1}{2t}$,*

$$\mathbb{E} \exp\left(\lambda \sup_{s \in [0, t]} \|B_s\|^2\right) \leq \left(\frac{1 + 2\lambda t}{1 - 2\lambda t}\right)^d.$$

Sampling guarantees.

We now show how to combine the continuous-time Rényi analysis ([Theorem 2.2.25](#)) with the discretization bound ([Theorem 6.2.1](#)) to obtain sampling guarantees. The first ingredient we need is a weak triangle inequality ([Exercise 6.5](#)).

Lemma 6.2.5 (Weak triangle inequality). *For any $q > 1$, $\lambda \in (0, 1)$, and any probability measures μ, ν, π :*

$$\mathcal{R}_q(\mu \parallel \pi) \leq \frac{q - \lambda}{q - 1} \mathcal{R}_{q/\lambda}(\mu \parallel \nu) + \mathcal{R}_{(q-\lambda)/(1-\lambda)}(\nu \parallel \pi). \quad (6.2.6)$$

In particular, for $\lambda = 1/2$,

$$\mathcal{R}_q(\mu \parallel \pi) \leq \frac{q - 1/2}{q - 1} \mathcal{R}_{2q}(\mu \parallel \nu) + \mathcal{R}_{2q-1}(\nu \parallel \pi).$$

Next, in order to apply [Theorem 6.2.1](#), we also need a concentration inequality for $\|\nabla V\|$ under π . The following result is taken from [Negrea \(2022\)](#); [Altschuler and Chewi \(2026a\)](#).

Lemma 6.2.7 (Score concentration). *Let $\pi \propto \exp(-V)$ and assume that ∇V is β -Lipschitz. Then, ∇V is $\sqrt{\beta}$ -sub-Gaussian under π , in the sense that for any $v \in \mathbb{R}^d$,*

$$\mathbb{E}_\pi \exp \langle \nabla V, v \rangle \leq \exp\left(\frac{\beta}{2} \|v\|^2\right).$$

In particular, for any $0 < \delta < 1/2$, with probability at least $1 - \delta$ under π ,

$$\|\nabla V\| \lesssim \sqrt{\beta d} + \sqrt{\beta \log \frac{1}{\delta}}.$$

Proof By β -smoothness,

$$\int \exp(-V(\cdot - v)) \geq \int \exp(-V + \langle \nabla V, v \rangle - \frac{\beta}{2} \|v\|^2).$$

Using $\int \exp(-V(\cdot - v)) = \int \exp(-V)$ and rearranging,

$$\exp\left(\frac{\beta}{2} \|v\|^2\right) \geq \int \exp \langle \nabla V, v \rangle \frac{\exp(-V)}{\int \exp(-V)} = \mathbb{E}_\pi \exp \langle \nabla V, v \rangle. \quad \square$$

We also need to transfer this concentration to different measures, which is accomplished via the following general principle.

Lemma 6.2.8 (Change of measure). *Suppose that a test function ϕ satisfies*

$$\pi\{\phi \geq \eta\} \leq \psi(\eta), \quad \text{for all } \eta > 0.$$

Then, for any measure μ and any $q > 1$,

$$\mu\{\phi \geq \eta\} \leq \psi(\eta)^{(q-1)/q} \exp\left(\frac{q-1}{q} \mathcal{R}_q(\mu \parallel \pi)\right).$$

Proof We note the following simple calculation:

$$\mu\{\phi \geq \eta\} = \int \mathbb{1}\{\phi \geq \eta\} \frac{d\mu}{d\pi} d\pi \leq \pi\{\phi \geq \eta\}^{(q-1)/q} \left\{ \int \left(\frac{d\mu}{d\pi}\right)^q d\pi \right\}^{1/q}. \quad \square$$

We can now prove the following result.

Theorem 6.2.9. *For $n \in \mathbb{N}$, let $\hat{\mu}_{nh}$ denote the law of the n -th iterate of LMC with step size $h > 0$. Assume that ∇V is β -Lipschitz, $q \geq 2$, and that ε is sufficiently small.*

1 *If π satisfies a log-Sobolev inequality, $C_{\text{LSI}}(\pi) \leq 1/\alpha$, then for an appropriate choice of h we obtain $\mathcal{R}_q(\hat{\mu}_{Nh} \parallel \pi) \leq \varepsilon^2$ with*

$$N = \Theta\left(\frac{\kappa^2 dq^3}{\varepsilon^2} \log^2 \frac{\mathcal{R}_{2q}(\hat{\mu}_0 \parallel \pi)}{\varepsilon^2}\right) \quad \text{iterations}.$$

2 *If π satisfies a Poincaré inequality, $C_{\text{PI}}(\pi) \leq 1/\alpha$, then for an appropriate choice of h we obtain $\mathcal{R}_q(\hat{\mu}_{Nh} \parallel \pi) \leq \varepsilon^2$ with*

$$N = \Theta\left(\frac{\kappa^2 dq^3}{\varepsilon^2} \left(\mathcal{R}_{2q}(\hat{\mu}_0 \parallel \pi)^2 + \log^2 \frac{1}{\varepsilon}\right)\right) \quad \text{iterations}.$$

Proof Let $(\pi_t)_{t \geq 0}$ denote the Langevin diffusion started at $\hat{\mu}_0$. By the weak triangle inequality (Lemma 6.2.5),

$$\mathcal{R}_q(\hat{\mu}_T \parallel \pi) \leq \mathcal{R}_{2q}(\hat{\mu}_T \parallel \pi_T) + \mathcal{R}_{2q-1}(\pi_T \parallel \pi).$$

By Theorem 2.2.25, the second term is at most ε^2 for $T \asymp (q/\alpha) \log(\mathcal{R}_{2q-1}(\hat{\mu}_0 \parallel \pi)/\varepsilon^2)$ in the LSI case, and for $T \asymp (q/\alpha) (\mathcal{R}_{2q-1}(\hat{\mu}_0 \parallel \pi) + \log 1/\varepsilon)$ in the PI case.

Next, Lemma 6.2.7 yields, for some universal $C > 0$,

$$\pi\{\|\nabla V\| \geq C\sqrt{\beta}(\sqrt{d} + \eta)\} \leq \exp(-\eta^2) \quad \text{for all } \eta > 0.$$

By Lemma 6.2.8, for any $t \in [0, T]$,

$$\pi_t\{\|\nabla V\| \geq C\sqrt{\beta}(\sqrt{d} + \eta)\} \leq \exp\left(-\frac{2q-1}{2q}(\eta^2 - \mathcal{R}_{2q}(\pi_t \parallel \pi))\right).$$

By the data-processing inequality (Theorem 1.5.6), $\mathcal{R}_{2q}(\pi_t \parallel \pi) \leq \mathcal{R}_{2q}(\hat{\mu}_0 \parallel \pi)$. We can conclude that with probability at least $1 - \delta$ under π_t ,

$$\|\nabla V\| \leq \sqrt{\beta(d + \mathcal{R}_{2q}(\hat{\mu}_0 \parallel \pi))} + \sqrt{\beta \log \frac{1}{\delta}}.$$

From Theorem 6.2.1, provided that $h \lesssim 1/\beta^2 q^2 T$ and $h \lesssim 1/\beta^{3/2} q T^{1/2}$,

$$\mathcal{R}_{2q}(\hat{\mu}_T \parallel \pi_T) \leq \beta^3 h^2 q T (d + \mathcal{R}_{2q}(\hat{\mu}_0 \parallel \pi)) + \beta^2 dhqT.$$

For simplicity, we use the assumption that ε is sufficiently small in order to focus on the dominant term (i.e., $\beta^2 dhqT$). The theorem follows by choosing h appropriately. $\square$

If we assume that $\mathcal{R}_{2q}(\hat{\mu}_0 \parallel \pi) = \tilde{O}(d)$ (see Exercise 6.7 for justification) and that $\kappa = O(1)$, then the complexity reads $\tilde{O}(dq^3/\varepsilon^2)$ in the LSI case and $\tilde{O}(d^3 q^3/\varepsilon^2)$ in the PI case. In the LSI case, this is worse than the result in Theorem 6.1.2 in two ways: first, the dependence on q (because here we did not take advantage of hypercontractivity); and second, Theorem 6.2.9 does not allow for taking an

unbounded number of steps of **LMC** (it suggests that if we run **LMC** for too long, then we will start to drift farther away from π , which is absurd).

As mentioned above, however, the benefit of the Girsanov approach is its flexibility, which allowed us to tackle the Poincaré case. This flexibility will be important when we study the underdamped Langevin diffusion in the next section. Finally, we remark that the sampling guarantee under a Poincaré inequality will be considerably improved via the proximal sampler in Chapter 8.

6.3 Analysis of ULMC via Girsanov's Theorem

In this section, we apply the Girsanov approach to provide Rényi divergence guarantees for **ULMC**, which was introduced in Section 5.3. This result provides a “warm start” for MALA, as discussed in Chapter 7.

Let $\mathbf{W}_T$ be the path measure under which B is a standard Brownian motion that drives the underdamped Langevin diffusion:

$$\begin{aligned} dX_t &= P_t dt, \\ dP_t &= -\gamma P_t dt - \nabla V(X_t) dt + \sqrt{2\gamma} dB_t. \end{aligned}$$

Then, let $\mathbf{P}_T$ be the path measure under which $\tilde{B}$ is a standard Brownian motion, which drives the interpolated **ULMC** algorithm:

$$\begin{aligned} dX_t &= P_t dt, \\ dP_t &= -\gamma P_t dt - \nabla V(X_{t_-}) dt + \sqrt{2\gamma} d\tilde{B}_t. \end{aligned}$$

We will establish the following discretization bound.

Theorem 6.3.1 (ULMC Rényi discretization bound). *Assume that the potential V is β -smooth and that the step size satisfies $h \lesssim 1/\sqrt{\beta} \wedge 1/\gamma \wedge 1/\beta^{2/3} q^{2/3} T^{1/3}$, where $T := Nh$. Then, for all $q \geq 2$,*

$$\begin{aligned} \mathcal{R}_q(\mathbf{P}_T \parallel \mathbf{W}_T) &\lesssim \frac{1}{q} \max_{n=0,1,\dots,N-1} \log \mathbb{E}^{\mathbf{W}_T} \exp \left\{ O \left(\frac{\beta^2 q^2 T}{\gamma} (h^2 \|P_{nh}\|^2 + h^4 \|\nabla V(X_{nh})\|^2) \right) \right\} \\ &\quad + \beta^2 d h^3 q T. \end{aligned}$$

Proof Since the argument is similar to the proof of [Theorem 6.2.1](#), we omit some of the steps for brevity. First, from Girsanov's theorem ([Theorem 3.2.8](#)), we obtain

$$\begin{aligned} \mathbb{E}^{\mathbf{W}_T} \left[\left(\frac{d\mathbf{P}_T}{d\mathbf{W}_T} \right)^q \right] &= \mathbb{E}^{\mathbf{W}_T} \exp \left(\frac{q}{\sqrt{2\gamma}} \int_0^T \langle \nabla V(X_{t_-}) - \nabla V(X_t), dB_t \rangle \right. \\ &\quad \left. - \frac{q}{4\gamma} \int_0^T \|\nabla V(X_{t_-}) - \nabla V(X_t)\|^2 dt \right). \end{aligned}$$

By the same martingale and AM–GM argument as before, it implies

$$\begin{aligned} \mathbb{E}^{\mathbf{W}_T} \left[\left(\frac{d\mathbf{P}_T}{d\mathbf{W}_T} \right)^q \right] &\leq \mathbb{E}^{\mathbf{W}_T} \exp \left(\frac{\beta^2 q^2}{\gamma} \int_0^T \|X_t - X_{t_-}\|^2 dt \right) \\ &\leq \frac{1}{T} \int_0^T \mathbb{E}^{\mathbf{W}_T} \exp \left(\frac{\beta^2 q^2 T}{\gamma} \|X_t - X_{t_-}\|^2 \right) dt. \end{aligned}$$

Next, we need an almost-sure version of the movement bound in [Lemma 5.3.11](#); we leave it for the reader to check that if $h \lesssim 1/\sqrt{\beta} \wedge 1/\gamma$, then

$$\|X_t - X_{t-}\| \lesssim h \|P_{t-}\| + h^2 \|\nabla V(X_{t-})\| + \gamma^{1/2} h \sup_{s \in [t-, t-+h]} \|B_s - B_{t-}\|.$$

Using the Brownian tail bound from [Lemma 6.2.4](#), it yields

$$\begin{aligned} \mathbb{E}^{\mathbf{w}_T} \exp\left(\frac{\beta^2 q^2 T}{\gamma} \|X_t - X_{t-}\|^2\right) &\leq \mathbb{E}^{\mathbf{w}_T} \exp\left(O\left(\frac{\beta^2 q^2 T}{\gamma} (h^2 \|P_{t-}\|^2 + h^4 \|\nabla V(X_{t-})\|^2)\right)\right) \\ &\quad \times \exp(O(\beta^2 d h^3 q^2 T)) \end{aligned}$$

provided that $h^3 \lesssim 1/\beta^2 q^2 T$. This completes the proof. $\square$

As in the previous section, this discretization bound leads to sampling guarantees, which we leave as an exercise.

Corollary 6.3.2 (ULMC guarantees). *For $n \in \mathbb{N}$, let $\hat{\mu}_{nh}$ denote the law of $\hat{X}_{nh}$ in [ULMC](#) with step size $h > 0$. Assume that ∇V is β -Lipschitz and for simplicity let us assume that $\mathcal{R}_2(\hat{\mu}_0 \parallel \pi) = \tilde{O}(d)$. The following guarantees hold with appropriate choices of h and γ .*

1 *If π satisfies a log-Sobolev inequality, $C_{\text{LSI}}(\pi) \leq 1/\alpha$, then $\|\hat{\mu}_{Nh} - \pi\|_{\text{TV}} \leq \varepsilon$ with*

$$N = \tilde{\Theta}\left(\frac{\kappa^{3/2} d^{1/2}}{\varepsilon}\right) \quad \text{iterations}.$$

If, moreover, π is log-concave, then for all $q \geq 2$, we can achieve $\mathcal{R}_q(\hat{\mu}_{Nh} \parallel \pi) \leq \varepsilon^2$ with

$$N = \tilde{\Theta}\left(\frac{\kappa d^{1/2} q^{1/2}}{\varepsilon}\right) \quad \text{iterations}.$$

2 *If π satisfies a Poincaré inequality, $C_{\text{PI}}(\pi) \leq 1/\alpha$, then we obtain $\mathcal{R}_{3/2}(\hat{\mu}_{Nh} \parallel \pi) \leq \varepsilon^2$ with*

$$N = \tilde{\Theta}\left(\frac{\kappa^{3/2} d^2}{\varepsilon}\right) \quad \text{or} \quad \tilde{\Theta}\left(\frac{\kappa d^2}{\varepsilon}\right) \quad \text{iterations},$$

where the second rate holds when π is log-concave.

Despite the fact that [Theorem 6.3.1](#) holds for all Rényi orders $q \geq 2$, the stated guarantees hold in weaker metrics as we need to weaken the Rényi order when we apply the weak triangle inequality.

In [Exercise 6.9](#), we outline an alternative approach to establishing a Rényi divergence guarantee for [ULMC](#) in the strongly log-concave case which bypasses hypocoercivity ([Theorem 5.3.30](#)).

Bibliographical Notes

The survey [van Erven and Harremoës \(2014\)](#) is a good general reference for Rényi divergences.

Rényi guarantees for [LMC](#) were first considered in [Vempala and Wibisono \(2019\)](#), which proved convergence of [LMC](#) to its biased stationary distribution provided that the biased limit satisfies a Poincaré or log-Sobolev inequality. However, this does not lead to a sampling guarantee unless the size of the “Rényi bias” (the Rényi divergence between the biased stationary distribution and the true target distribution) can be estimated. Mixing time bounds toward the biased stationary distribution were also later considered in [Altschuler and Talwar \(2023\)](#).

Motivated by applications to differential privacy, [Ganesh and Talwar \(2020\)](#) provided the first Rényi sampling guarantees for LMC under strong log-concavity by using a technique based on the adaptive composition lemma for Rényi divergences. This result was refined in [Erdogdu et al. \(2022\)](#) via a two-phase analysis, still relying on the adaptive composition lemma. Subsequently, building on the earlier work of [Chewi et al. \(2021\)](#) (which essentially contains a one-step Rényi discretization argument), it was realized in [Chewi et al. \(2025a\)](#) that the earlier arguments of [Ganesh and Talwar \(2020\)](#); [Erdogdu et al. \(2022\)](#) can be streamlined by replacing the adaptive composition lemma entirely with Girsanov's theorem. The proofs of [Theorem 6.1.2](#), [Exercise 6.2](#), and [Exercise 6.7](#) are also from [Chewi et al. \(2025a\)](#). The proof of [Theorem 6.2.9](#) given here is a further refinement of [Chewi et al. \(2025a\)](#) using the score concentration lemma ([Lemma 6.2.7](#)).

The argument of [Exercise 6.3](#) is inspired by [Boucheron et al. \(2013, Exercise 5.16\)](#) and the Rényi bounds were recorded in [Liu and Chewi \(2026\)](#). The second part is the standard Herbst argument.

The use of Girsanov's theorem to obtain Rényi discretization bounds for ULMC first appeared in [Zhang et al. \(2023\)](#), which considered a variety of settings (including weakly smooth potentials); that paper also established the first KL divergence guarantee with $\sqrt{d}$ dimension dependence for ULMC in the strongly log-concave and log-smooth case. The Girsanov argument was substantially simplified in [Altschuler and Chewi \(2024a\)](#).

The first Rényi divergence guarantee for ULMC with $\sqrt{d}$ dimension dependence was established in [Altschuler and Chewi \(2024a\)](#) via the approach outlined in [Exercise 6.9](#). This approach is closely related to the Harnack inequality, which we have already encountered for the Langevin diffusion in [Exercise 2.17](#) and [Exercise 3.4](#); see the Bibliographical Notes to Chapter 2. The corresponding inequality for the underdamped Langevin diffusion is rather more difficult to establish. Whereas [Exercise 6.9](#) establishes a regularizing inequality for one step of the discrete-time algorithm, a full Harnack inequality for the diffusion was first shown in [Guillin and Wang \(2012, Corollary 4.6\)](#); see also [Altschuler et al. \(2025, Theorem 3.2\)](#), where it is established using the method of [Exercise 3.4](#).

Exercises

Analysis of LMC via Interpolation Argument

▷ [Exercise 6.1](#) (Rényi differential inequality)

Prove [Proposition 6.1.1](#) and [Proposition 6.1.4](#).

▷ [Exercise 6.2](#) (Rényi discretization bound for log-concave targets)

Suppose that $\pi \propto \exp(-V)$ is log-concave, that $\nabla V(0) = 0$, and that ∇V is β -Lipschitz. Also, suppose that we initialize LMC at $\hat{\mu}_0 = \text{normal}(0, \beta^{-1}I_d)$. The goal of this exercise is to prove a Rényi discretization bound for LMC under these assumptions.

- 1 First, show that $\hat{\mu}_t$ satisfies an LSI for all $t \geq 0$, and write down a bound for $C_{\text{LSI}}(\hat{\mu}_t)$ (the bound should grow linearly with t).

Hint: See [Exercise 4.4](#).

- 2 Follow the proof of [Theorem 6.1.2](#) (the first proof which incurs a suboptimal dependence on q). Note that in [Theorem 6.1.2](#), we bounded $\text{KL}(\mathbf{P} \parallel \mathbb{P}) \leq \frac{q-1}{q} \text{KL}(\psi_t \hat{\mu}_t \parallel \pi)$ and we applied the LSI for π . This time, use $\text{KL}(\mathbf{P} \parallel \mathbb{P}) = \text{KL}(\psi_t \hat{\mu}_t \parallel \hat{\mu}_t)$ and apply the LSI for $\hat{\mu}_t$ instead.

Also, instead of using the decay of the Rényi divergence under a LSI, use the decay of the Rényi divergence under a PI ([Theorem 2.2.25](#)). (Since π is log-concave, it necessarily satisfies a Poincaré inequality with some constant $1/\alpha$, see the Bibliographical Notes to Chapter 2.)

Prove that if $\varepsilon \leq \sqrt{1/q} \wedge \sqrt{\kappa d}$, then with an appropriate choice of step size h and with $N = \tilde{\Theta}(\kappa^2 d^2 q^3 / \varepsilon^2)$ iterations of **LMC**, we obtain $\sqrt{\mathcal{R}_q(\hat{\mu}_{Nh} \parallel \pi)} \leq \varepsilon$. (Unlike the guarantee of **Theorem 6.1.2**, here the guarantee does not allow N to be too large, due to the growing LSI constant of the iterates.)

▷ **Exercise 6.3** (Rényi analogues of the LSI)

Let π satisfy an LSI with constant $1/\alpha$ and let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be continuously differentiable.

- 1 Show that for any $0 < \lambda^2 < 2\alpha$,

$$\frac{\text{ent}_{\pi} \exp(\lambda f)}{\mathbb{E}_{\pi} \exp(\lambda f)} \leq \frac{\lambda^2}{2\alpha - \lambda^2} \log \mathbb{E}_{\pi} \exp(\|\nabla f\|^2).$$

Hint: Start by applying the log-Sobolev inequality to $\text{ent}_{\pi} \exp(\lambda f)$, and then apply the Donsker–Varadhan inequality with the measure ν given by $\frac{d\nu}{d\pi} = \frac{\exp(\lambda f)}{\mathbb{E}_{\pi} \exp(\lambda f)}$.

- 2 Let $\psi(\lambda) := \log \mathbb{E}_{\pi} \exp(\lambda (f - \mathbb{E}_{\pi} f))$. Show that

$$\frac{\text{ent}_{\pi} \exp(\lambda f)}{\mathbb{E}_{\pi} \exp(\lambda f)} = \lambda \psi'(\lambda) - \psi(\lambda)$$

and hence deduce that

$$\partial_{\lambda} \left(\frac{\psi(\lambda)}{\lambda} \right) \leq \frac{1}{2\alpha - \lambda^2} \log \mathbb{E}_{\pi} \exp(\|\nabla f\|^2).$$

Integrate this inequality to prove that

$$\log \mathbb{E}_{\pi} \exp(\lambda (f - \mathbb{E}_{\pi} f)) \leq \frac{\lambda^2}{2\alpha - \lambda^2} \log \mathbb{E}_{\pi} \exp(\|\nabla f\|^2).$$

Note that a Poincaré inequality bounds the variance of f in terms of the second moment of $\|\nabla f\|$; here, we have shown that a log-Sobolev inequality bounds the moment-generating function of f in terms of a sub-Gaussian moment of $\|\nabla f\|$.

- 3 Show that the preceding inequality implies that for any probability measure $\mu \ll \pi$ and $\lambda > q^2/2\alpha$,

$$\mathcal{R}_q(\mu \parallel \pi) \leq \frac{q^2}{2\alpha\lambda - q^2} \log \mathbb{E}_{\pi} \exp(\lambda \|\nabla \log \frac{\mu}{\pi}\|^2).$$

- 4 In fact, the expectation on the right-hand side can also be evaluated under μ if it is more convenient: for $\lambda > q(q-1)/2\alpha$,

$$\mathcal{R}_q(\mu \parallel \pi) \leq \frac{q^2}{2\alpha\lambda - q(q-1)} \log \mathbb{E}_{\mu} \exp(\lambda \|\nabla \log \frac{\mu}{\pi}\|^2).$$

Prove this by repeating the argument of the first part, but use Donsker–Varadhan to change the measure to μ (instead of π).

Analysis of LMC via Girsanov's Theorem

▷ **Exercise 6.4** (Sub-Gaussian bound for Brownian motion)

Prove **Lemma 6.2.4**.

Hint: The reflection principle states that if $(\tilde{B}_t)_{t \geq 0}$ is a one-dimensional Brownian motion, then for every $\eta > 0$, $\mathbb{P}(\sup_{s \in [0, t]} \tilde{B}_s > \eta) = 2 \mathbb{P}(\tilde{B}_t > \eta)$.

▷ **Exercise 6.5** (Weak triangle inequality)

Prove the weak triangle inequality (Lemma 6.2.5). What happens if we take $q = 1 + \varepsilon$, $\lambda = 1 - \varepsilon$, and send $\varepsilon \searrow 0$?

▷ **Exercise 6.6** (Score concentration in the non-smooth case)

Extend Lemma 6.2.7 to the case where ∇V is simply assumed to be Hölder continuous with some exponent $s \in [0, 1]$:

$$\|\nabla V(x) - \nabla V(y)\| \leq \beta \|x - y\|^s, \quad \text{for all } x, y \in \mathbb{R}^d.$$

▷ **Exercise 6.7** (A more general initialization bound)

Suppose that V is β -smooth, but not necessarily convex. In this case, it is not realistic to assume that we have access to the minimizer of V , but we can at least assume that we have computed a stationary point: $\nabla V(0) = 0$. Show that for $\hat{\mu}_0 = \text{normal}(0, (2\beta)^{-1}I_d)$,

$$\mathcal{R}_\infty(\hat{\mu}_0 \parallel \pi) \leq 2 + V(0) - \inf V + \frac{d}{2} \log(4\beta m^2),$$

where $m := \mathbb{E}_\pi \|\cdot\|$ is the first moment of π . This shows that if we can compute a stationary point with suboptimality gap $\tilde{O}(d)$, then it is reasonable to assume $\mathcal{R}_\infty(\hat{\mu}_0 \parallel \pi) = \tilde{O}(d)$.

Hint: Use the decomposition

$$\sup \frac{\hat{\mu}_0}{\pi} \leq \sup_{x \in \mathbb{R}^d} \exp(V(x) - \beta \|x\|^2) \frac{\int \exp(-V) \int \exp(-V - \delta \|\cdot\|^2)}{\int \exp(-V - \delta \|\cdot\|^2) (\pi/\beta)^{d/2}}$$

and bound each ratio separately.

Analysis of ULMC via Girsanov's Theorem

▷ **Exercise 6.8** (Rényi guarantees for ULMC)

This question fills some omitted details from the proofs.

- 1 Fill in the missing details in the proof of Theorem 6.3.1.
- 2 Prove Corollary 6.3.2 by combining Theorem 6.3.1 with the hypocoercivity results from Section 5.3.

▷ **Exercise 6.9** (Discrete-time version of hypocoercive regularization)

In this exercise, we outline a simpler, discrete-time alternative to Rényi convergence of the underdamped Langevin diffusion which bypasses hypocoercivity (Theorem 5.3.30).

Let Σ denote the covariance matrix for a single step of ULMC, as given in Lemma 5.3.9. Recall, however, that for the contraction result we work with the twisted coordinates $(x, x + \frac{2}{\gamma} p)$. Let $\mathcal{M}$ denote the matrix corresponding to the change of coordinates $(x, p) \mapsto (x, x + \frac{2}{\gamma} p)$, so that the covariance matrix in this system of coordinates is $\tilde{\Sigma} := \mathcal{M}\Sigma\mathcal{M}^\top$. Also, let $\hat{\mathbf{P}}$ denote the Markov kernel corresponding to one iteration of ULMC.

- 1 Prove that for $h \lesssim 1/\gamma$, $\lambda_{\min}(\tilde{\Sigma}) \gtrsim \gamma h^3$.
- 2 Using the explicit expression for the Rényi divergence between Gaussians, prove the following regularization bound:

$$\mathcal{R}_q(\delta_{x,p} \hat{\mathbf{P}} \parallel \delta_{\bar{x},\bar{p}} \hat{\mathbf{P}}) \lesssim \frac{q}{\gamma h^3} \|(x, p) - (\bar{x}, \bar{p})\|^2.$$

- 3 Although we proved that the diffusion contracts in this coordinate system ([Theorem 5.3.6](#)), our aim now is to work with the discrete-time algorithm. Adapt the proof of [Theorem 5.3.6](#), using [Lemma 5.3.9](#), to show that if $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and $h \lesssim 1/\beta^{1/2}\kappa$, then $\hat{\mathbf{P}}$ is an almost-sure contraction:

$$W_{\infty, \|\cdot\|}(\delta_{x,p} \hat{\mathbf{P}}, \delta_{\bar{x}, \bar{p}} \hat{\mathbf{P}}) \leq \left(1 - \Omega\left(\frac{\alpha h}{\sqrt{\beta}}\right)\right) \|\!(x, p) - (\bar{x}, \bar{p})\!\|,$$

where $W_{\infty, \|\cdot\|}$ denotes the W_{∞} metric with respect to $\|\cdot\|$.

- 4 Combine these facts to prove that $\mathcal{R}_q(\hat{\mu}_0 \hat{\mathbf{P}}^N \parallel \pi \hat{\mathbf{P}}^N) \leq \varepsilon^2$, where $\hat{\mu}_0 = \text{normal}(x_*, \beta^{-1} I_d) \otimes \text{normal}(0, I_d)$, provided that

$$N \gtrsim \frac{\sqrt{\beta}}{\alpha h} \log \frac{\kappa d q}{\beta^{3/2} \varepsilon^2 h^3}.$$

- 5 Combine this with [Theorem 6.3.1](#) to give another $\mathcal{R}_q$ guarantee for [ULMC](#) with iteration complexity $\tilde{O}(\kappa^{3/2} d^{1/2} q^{1/2} / \varepsilon)$.

≥ **Exercise 6.10** (Rényi guarantees in the log-concave case)

In [Altschuler et al. \(2025\)](#), the following bound was established. Let $(\mathbf{P}_t)_{t \geq 0}$ denote the Markov semigroup for the underdamped Langevin diffusion, and assume that $0 \leq \nabla^2 V \leq \beta I_d$ and $\gamma \geq \sqrt{32\beta}$. Then, for all $q \geq 1$, $T > 0$, and $x, \bar{x}, p, \bar{p} \in \mathbb{R}^d$,

$$\mathcal{R}_q(\delta_{x,p} \mathbf{P}_T \parallel \delta_{\bar{x}, \bar{p}} \mathbf{P}_T) \lesssim q \left(\frac{1}{\gamma T^3} + \frac{\gamma}{T} \right) \left(\|x - \bar{x}\|^2 + \frac{1}{\gamma^2 + 1/T^2} \|p - \bar{p}\|^2 \right).$$

Can you combine this with [Theorem 6.3.1](#) to establish Rényi guarantees for [ULMC](#) in this case?

So far, we have focused on discretizations of diffusions. Discretization of a continuous-time Markov process yields a discrete-time Markov chain whose stationary distribution is no longer equal to the target π ; we say that the algorithm is *biased*. Nevertheless, we showed that the size of the bias can be made smaller than any desired accuracy ε by choosing a small step size h , which then leads to quantitative sampling guarantees.

However, the number of iterations of the algorithm is proportional to the inverse step size $1/h$, and consequently the complexity of the algorithms scaled as $\text{poly}(1/\varepsilon)$. In this section, we address the problem of designing *high-accuracy* samplers, i.e., samplers whose complexity scales as $\text{polylog}(1/\varepsilon)$. To accomplish this, we must fix the bias of the sampling algorithm, which is accomplished via the Metropolis–Hastings filter.

7.1 Rejection Sampling

Before introducing the Metropolis–Hastings filter, we begin with a warm-up and introduce the concept of rejection via the **rejection sampling** algorithm.

Rejection sampling: Let π be the target distribution and let $\tilde{\pi}$ be an unnormalized version of π , i.e., $\tilde{\pi} \propto \pi$. Suppose we can sample from a distribution μ and that an unnormalized version $\tilde{\mu}$ of μ satisfies $\tilde{\mu} \geq \tilde{\pi}$. Then, repeat until acceptance:

- 1 Draw $X \sim \mu$.
- 2 Accept X with probability $\tilde{\pi}(X)/\tilde{\mu}(X)$.

Unlike the other sampling algorithms we have considered thus far, rejection sampling always terminates with an *exact* sample from π .

Theorem 7.1.1 (Rejection sampling). *The output of rejection sampling is a sample drawn exactly from π . Also, the number of samples drawn from μ until a sample is accepted follows a geometric distribution with mean Z_μ/Z_π , where $Z_\mu := \int \tilde{\mu}$ and $Z_\pi := \int \tilde{\pi}$.*

Proof The probability of acceptance is

$$\mathbb{P}(\text{acceptance}) = \int \frac{\tilde{\pi}}{\tilde{\mu}} d\mu = \frac{Z_\pi}{Z_\mu} \int \frac{\pi}{\mu} d\mu = \frac{Z_\pi}{Z_\mu}$$

and clearly the number of samples drawn until acceptance is geometrically distributed.

To show that the output X of rejection sampling is drawn exactly according to π , let $(U_i)_{i=1}^\infty \stackrel{\text{i.i.d.}}{\sim} \text{uniform}([0, 1])$ and $(X_i)_{i=1}^\infty \stackrel{\text{i.i.d.}}{\sim} \mu$ be independent. Then, for any event A ,

$$\begin{aligned} \mathbb{P}(X \in A) &= \sum_{n=0}^{\infty} \mathbb{P}\left(X_{n+1} \in A, U_i > \frac{\tilde{\pi}(X_i)}{\tilde{\mu}(X_i)} \forall i \in [n], U_{n+1} \leq \frac{\tilde{\pi}(X_{n+1})}{\tilde{\mu}(X_{n+1})}\right) \\ &= \sum_{n=0}^{\infty} \mathbb{P}\left(X_{n+1} \in A, U_{n+1} \leq \frac{\tilde{\pi}(X_{n+1})}{\tilde{\mu}(X_{n+1})}\right) \mathbb{P}\left(U_1 > \frac{\tilde{\pi}(X_1)}{\tilde{\mu}(X_1)}\right)^n \\ &= \sum_{n=0}^{\infty} \left(\int_A \frac{\tilde{\pi}}{\tilde{\mu}} d\mu\right) \left(\int \left(1 - \frac{\tilde{\pi}}{\tilde{\mu}}\right) d\mu\right)^n = \frac{Z_\pi}{Z_\mu} \pi(A) \sum_{n=0}^{\infty} \left(1 - \frac{Z_\pi}{Z_\mu}\right)^n = \pi(A). \quad \square \end{aligned}$$

The rejection sampling algorithm requires the construction of the upper envelope $\tilde{\mu}$. We now demonstrate how to construct this envelope for our usual class of distributions. Namely, suppose that $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$, and that $\nabla V(0) = 0$. We can assume that our unnormalized version $\tilde{\pi} = \exp(-V)$ of π satisfies $V(0) = 0$ (if not, replace V by $V - V(0)$). Then, by strong convexity of V , we see that $\tilde{\pi} \leq \exp(-\frac{\alpha}{2} \|\cdot\|^2)$, and we can take $\tilde{\mu} := \exp(-\frac{\beta}{2} \|\cdot\|^2)$, which means that the normalized distribution is $\mu = \text{normal}(0, \alpha^{-1} I_d)$. To understand the efficiency of rejection sampling, we need to bound the ratio Z_μ/Z_π of normalizing constants. By smoothness of V ,

$$\frac{Z_\mu}{Z_\pi} = \frac{(2\pi/\alpha)^{d/2}}{\int \exp(-V)} \leq \frac{(2\pi/\alpha)^{d/2}}{\int \exp(-\frac{\beta}{2} \|\cdot\|^2)} = \frac{(2\pi/\alpha)^{d/2}}{(2\pi/\beta)^{d/2}} = \kappa^{d/2},$$

with $\kappa := \beta/\alpha$. We summarize this result in the following proposition.

Proposition 7.1.2 (Rejection sampling under strong log-concavity). *Let the target $\pi \propto \exp(-V)$ on $\mathbb{R}^d$ satisfy $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$, $V(0) = 0$, and $\nabla V(0) = 0$. Then, rejection sampling with the envelope $\tilde{\mu} := \exp(-\frac{\beta}{2} \|\cdot\|^2)$ returns an exact sample from π with a number of iterations that is a geometric random variable with mean at most $\kappa^{d/2}$, where $\kappa := \beta/\alpha$.*

The rejection sampling guarantee can be formulated in one of two ways. We can think of the algorithm as returning an exact sample from π , with a random number of iterations (the number of iterations is geometrically distributed). Alternatively, if we place an upper bound N on the number of iterations of the algorithm and output “FAIL” if we have not terminated by iteration N , then the probability of “FAIL” is at most $\varepsilon := (1 - 1/\kappa^{d/2})^N$, and if μ_N denotes the law of the output of the algorithm, then $\|\mu_N - \pi\|_{\text{TV}} \leq \varepsilon$. If we flip this around and fix the target accuracy ε , we see that the number of iterations required to achieve this guarantee is $N \geq \kappa^{d/2} \log(1/\varepsilon)$.

Although this result is acceptable in low dimension, the complexity of this approach quickly becomes intractable even for moderately high-dimensional problems. In the next section, we will see that by combining the idea of rejection with *local* proposals, we can obtain tractable sampling algorithms in high dimension.

7.2 The Metropolis–Hastings Filter

A Metropolis–Hastings algorithm consists of *proposing* moves from a proposal kernel Q , and then *accepting or rejecting* each move with a carefully chosen probability which ensures that the resulting Markov chain has the desired stationary distribution π .

In more detail, let Q be a kernel on $\mathbb{R}^d \times \mathbb{R}^d$, that is: for each $x \in \mathbb{R}^d$, $Q(x, \cdot)$ is a probability measure on $\mathbb{R}^d$. We will mostly consider proposals such that each $Q(x, \cdot)$ has a density with respect to Lebesgue measure, and via an abuse of notation we will write $Q(x, y)$ for this density evaluated at y (an exception is when we consider MHMC below).

Starting from $X \in \mathbb{R}^d$, we tentatively propose a new point $Y \sim Q(X, \cdot)$. We then accept the point Y with probability $A(X, Y)$ (called the *acceptance probability*); otherwise, we stay at the old point X . Iterate this process until convergence.

There are different possible choices for the acceptance probability A , but the choice we consider here is the **Metropolis–Hastings filter**

$$A(x, y) := 1 \wedge \frac{\pi(y) Q(y, x)}{\pi(x) Q(x, y)}. \quad (7.2.1)$$

The overall algorithm is summarized as follows.

Metropolis–Hastings algorithm (with proposal Q): initialize at a point $X_0 \in \mathbb{R}^d$. Then, iterate the following steps for $n = 1, 2, 3, \dots$:

- 1 Propose a new point $Y_n \sim Q(X_{n-1}, \cdot)$.
- 2 With probability $A(X_{n-1}, Y_n)$, set $X_n := Y_n$; otherwise, set $X_n := X_{n-1}$. Here, A is the acceptance probability defined via (7.2.1).

This algorithm defines a discrete-time Markov chain whose transition kernel P can be written explicitly as

$$P(x, dy) = Q(x, dy) A(x, y) + \underbrace{\left(1 - \int Q(x, dy') A(x, y')\right)}_{\text{rejection probability}} \delta_x(dy). \quad (7.2.2)$$

A discrete-time Markov chain with transition kernel P is called **reversible** with respect to π if it holds that $\pi(dx) P(x, dy) = \pi(dy) P(y, dx)$. Similarly to our discussion in Section 1.2, discrete-time reversible Markov chains can be studied via spectral theory.

Theorem 7.2.3 (Reversibility of Metropolis–Hastings). *The Metropolis–Hastings algorithm with proposal Q is reversible with respect to π .*

Proof We want to check that $\pi(x) P(x, y) = \pi(y) P(y, x)$ for all $x, y \in \mathbb{R}^d$ with $x \neq y$. We can write

$$\begin{aligned} \pi(x) P(x, y) &= \pi(x) Q(x, y) A(x, y) = \pi(x) Q(x, y) \min\left\{1, \frac{\pi(y) Q(y, x)}{\pi(x) Q(x, y)}\right\} \\ &= \min\{\pi(x) Q(x, y), \pi(y) Q(y, x)\} \end{aligned}$$

and this expression is symmetric in x and y . □

Take note of the flexibility of the Metropolis–Hastings algorithm! Regardless of the choice of proposal kernel Q , the filter always makes the algorithm unbiased. Of course, the choice of Q will be crucial later in order to guarantee rapid convergence to stationarity.

Implementability of the Metropolis–Hastings algorithm.

To implement the algorithm, the proposal Q must be simple enough such that (1) we can sample from $Q(x, \cdot)$ easily, and (2) we can compute the density $Q(x, y)$ easily (which is required to compute the acceptance probability). Note that although the target density π appears in the expression (7.2.1) for the acceptance probability, it only appears as a *ratio*, and in particular we do not need to know the normalization constant of π . Hence, the Metropolis–Hastings filter can be implemented using queries to the density of π up to normalization, which are “zeroth-order queries” (unlike, e.g., LMC, which uses first-order information through queries to the gradient ∇V).

Metropolis–Hastings as a projection.

There is a nice geometric interpretation of the Metropolis–Hastings filter as a projection, due to Billera and Diaconis (2001). Given a proposal kernel Q , let $P(Q)$ denote the Metropolis–Hastings kernel obtained from Q (see (7.2.2)). Then, the mapping $Q \mapsto P(Q)$ is a projection of the proposal kernel Q onto the space of reversible Markov chains with stationary distribution π with respect to an L^1 notion of distance, defined as follows:

$$d(P, P') := \int_{(\mathbb{R}^d \times \mathbb{R}^d) \setminus \text{diag}} |P(x, y) - P'(x, y)| \pi(\mathrm{d}x) \mathrm{d}y \quad (7.2.4)$$

where $\text{diag} := \{(x, x) \mid x \in \mathbb{R}^d\}$ is the diagonal in $\mathbb{R}^d \times \mathbb{R}^d$.

Theorem 7.2.5 (Billera and Diaconis (2001)). *Let $\mathcal{R}(\pi)$ denote the space of kernels P which are reversible with respect to π , and such that for each $x \in \mathbb{R}^d$, $P(x, \cdot)$ admits a density with respect to Lebesgue measure (except possibly having an atom at x). Then,*

$$P(Q) \in \arg \min_{P \in \mathcal{R}(\pi)} d(P, Q).$$

Proof Let $P \in \mathcal{R}(\pi)$, and let $S := \{(x, y) \in \mathbb{R}^d \times \mathbb{R}^d \mid \pi(x) Q(x, y) > \pi(y) Q(y, x)\}$. Then,

$$\begin{aligned} d(P, Q) &= \int_{(\mathbb{R}^d \times \mathbb{R}^d) \setminus \text{diag}} |P(x, y) - Q(x, y)| \pi(\mathrm{d}x) \mathrm{d}y \\ &= \int_S |P(x, y) - Q(x, y)| \pi(\mathrm{d}x) \mathrm{d}y + \int_{(\mathbb{R}^d \times \mathbb{R}^d) \setminus (S \cup \text{diag})} |P(x, y) - Q(x, y)| \pi(\mathrm{d}x) \mathrm{d}y \\ &= \int_S |P(x, y) - Q(x, y)| \pi(\mathrm{d}x) \mathrm{d}y + \int_S |P(y, x) - Q(y, x)| \pi(\mathrm{d}y) \mathrm{d}x. \end{aligned}$$

Using reversibility of P , the second term is

$$\begin{aligned} \int_S |P(y, x) - Q(y, x)| \pi(\mathrm{d}y) \mathrm{d}x &= \int_S |\pi(x) P(x, y) - \pi(y) Q(y, x)| \mathrm{d}x \mathrm{d}y \\ &\geq \int_S |\pi(x) Q(x, y) - \pi(y) Q(y, x)| \mathrm{d}x \mathrm{d}y - \int_S |P(x, y) - Q(x, y)| \pi(\mathrm{d}x) \mathrm{d}y. \end{aligned} \quad (7.2.6)$$

Putting this together, $d(P, Q) \geq \int_S |\pi(x) Q(x, y) - \pi(y) Q(y, x)| \mathrm{d}x \mathrm{d}y$, so we have obtained a lower bound which does not depend on P . On the other hand, we can check that the only inequality (7.2.6) that we used is an equality for $P = P(Q)$. $\square$

7.3 An Overview of High-Accuracy Samplers

As we already discussed, the Metropolis–Hastings framework is quite flexible: by instantiating it with different choices for the proposal Q , we obtain several different algorithms.

Metropolized random walk (MRW).

Perhaps the simplest proposal is to take $Q(x, \cdot) = \text{normal}(x, hI_d)$, which yields the **Metropolized random walk (MRW)** algorithm. This corresponds to simply taking a random walk around the state space, where some steps are occasionally rejected. Note that since the proposal is independent of the target π , the overall algorithm only uses queries to the density of π up to normalization (to implement the filter); thus, it is the only algorithm we have discussed so far (besides rejection sampling) which uses only a zeroth-order oracle for the potential V .

Metropolis-adjusted Langevin algorithm (MALA).

A better choice of proposal is

$$Q(x, \cdot) = \text{normal}(x - h \nabla V(x), 2h I_d)$$

which is one step of the **LMC** algorithm; this yields the **Metropolis-adjusted Langevin algorithm (MALA)**. We will carefully study the convergence guarantees for MALA in this chapter.

Metropolized Hamiltonian Monte Carlo (MHMC).

Recall the Hamiltonian Monte Carlo (HMC) algorithm that we introduced in Section 5.2. The ideal HMC algorithm is not implementable because it requires the ability to exactly integrate Hamilton's equations, and this is generally not possible outside of a few special cases.

We now consider approximately implementing Hamilton's equations through the use of a numerical integrator. Although several choices are available, for Hamilton's equations it is preferable to use a **symplectic integrator**.¹ We will focus on the simplest and most well-known such integrator, called the **leapfrog integrator**.

Leapfrog integrator: Pick a step size $h > 0$ and a total number of iterations K , corresponding to the total integration time via $T := Kh$. Let (x_0, p_0) be the initial point. For $k = 0, 1, 2, \dots, K - 1$:

- 1 Set $p_{(k+\frac{1}{2})h} := p_{kh} - \frac{h}{2} \nabla V(x_{kh})$.
- 2 Set $x_{(k+1)h} := x_{kh} + h p_{(k+\frac{1}{2})h}$.
- 3 Set $p_{(k+1)h} := p_{(k+\frac{1}{2})h} - \frac{h}{2} \nabla V(x_{(k+1)h})$.

Once we apply the leapfrog integrator to HMC, we obtain a discrete-time sampling algorithm which is once again biased. We then correct the bias through the use of the Metropolis–Hastings filter. Specifically, for an integration time $T = Kh$, let

$$F_{\text{leap}}(x, p) = \text{output } x_T \text{ of the leapfrog integrator with } K \text{ steps,} \\ \text{started at } (x, p) .$$

Remarkably, the acceptance probability can be computed in closed form, and this relies on specific properties of the leapfrog integrator. The full algorithm is summarized as follows.

¹When placed within the framework of geometry, Hamiltonian mechanics is encoded via *symplectic geometry*, which is the study of manifolds equipped with a symplectic 2-form. The flow map for Hamilton's equations preserves this symplectic form, and is therefore known as a *symplectomorphism*. Symplectic integrators are special integrators which also preserve the symplectic form. This property leads to stability, especially for long integration times.

Metropolized Hamiltonian Monte Carlo (MHMC): Initialize at $X_0 \sim \mu_0$. For iterations $n = 0, 1, 2, \dots$:

- 1 *Refresh* the momentum: draw $P_n \sim \text{normal}(0, I_d)$.
- 2 *Propose* a trajectory: let $(X'_n, P'_n) := F_{\text{leap}}(X_n, P_n)$.
- 3 *Accept* the trajectory with probability $1 \wedge \exp\{H(X_n, P_n) - H(X'_n, P'_n)\}$. If the trajectory is accepted, set $X_{n+1} := X'_n$; otherwise, we set $X_{n+1} := X_n$.

It turns out that when $K = 1$, the MHMC algorithm reduces to MALA ([Exercise 7.1](#)).

We next justify why the MHMC algorithm leaves π invariant. Actually, although we have written down the MHMC algorithm in the form which is easiest to implement, it obscures the underlying structure of the algorithm. The proof of the next theorem will clarify this point.

Theorem 7.3.1 (Correctness of MHMC). *The augmented target distribution $\pi \propto \exp(-H)$ is invariant for the MHMC algorithm.*

Proof First, we note that the step of refreshing the momentum leaves π invariant, so it suffices to study the proposal and acceptance steps.

For the moment, let us pretend that the proposal actually uses the true flow map F_T (which exactly integrates Hamilton's equations for time T) rather than the leapfrog integrator F_{leap} . Then, the proposal kernel is deterministic, $Q((x, p), \cdot) = \delta_{F_T(x, p)}$.²

If we naively apply the Metropolis–Hastings filter, the probability of accepting a proposal (x', p') starting from (x, p) involves a ratio $Q((x', p'), (x, p)) / Q((x, p), (x', p'))$, but this ratio is ill-defined in our setting. The problem is that if $(x', p') = F_T(x, p)$, then it is not the case that $(x, p) = F_T(x', p')$; the proposal is not reversible. Hence, we would be led to reject every single trajectory.

To fix this, recall from [Exercise 5.5](#) that we have the following time reversibility property: if R denotes the momentum flip operator $(x, p) \mapsto (x, -p)$, then it holds that $F_T^{-1} = R \circ F_T \circ R$. It implies that $F_T \circ R = R \circ F_T^{-1} = (F_T \circ R)^{-1}$, so $F_T \circ R$ is idempotent. In other words, if we use the proposal $F_T \circ R$ (i.e., first flip the momentum before integrating Hamilton's equations), then the proposal would be reversible and the above issue does not arise, as the ratio $Q((x', p'), (x, p)) / Q((x, p), (x', p'))$ would equal 1. Observe also that using $F_T \circ R$ instead of F_T does not change the algorithm since we refresh the momentum at each step (and if $P_n \sim \text{normal}(0, I_d)$, then $-P_n \sim \text{normal}(0, I_d)$ as well).

Once we use the proposal $(x', p') = (F_T \circ R)(x, p) = F_T(x, -p)$, the Metropolis–Hastings acceptance probability is calculated to be

$$1 \wedge \frac{\pi(x', p')}{\pi(x, p)} = 1 \wedge \exp\{H(x, p) - H(x', p')\}. \quad (7.3.2)$$

When we use the exact flow map F_T , then the Hamiltonian is conserved ([Exercise 5.5](#)) so the above probability is one; every trajectory is accepted. However, the above expression is indeed meaningful if we instead use the leapfrog integrator F_{leap} .

So far, we have motivated the expression (7.3.2) based on the exact flow map F_T , but clearly the above argument holds just as well for the leapfrog integrator F_{leap} as soon as we verify the property $F_{\text{leap}}^{-1} = R \circ F_{\text{leap}} \circ R$, and this is where we use the specific form of the leapfrog integrator. We leave the verification as [Exercise 7.2](#). $\square$

²Up until now, we have been assuming that the proposal kernel admits a density w.r.t. Lebesgue measure, which certainly does not hold here, but we will brush over this technicality as it is not the key point here.

Remark 7.3.3. The proof shows that the proposal of MHMC should really be thought of as $F_{\text{leap}} \circ R$, instead of F_{leap} . In fact, if we did not refresh the momentum, then repeatedly applying the idempotent operator $F_{\text{leap}} \circ R$ would just cause the algorithm to jump back and forth between two points (x, p) and (x', p') , which is silly; hence one should also apply another momentum flip after the filter. In symbols, if MH denotes the Metropolis–Hastings filter step, and Refresh denotes the momentum refreshment step, we should think of MHMC as the composition

$$\text{MHMC} = R \circ \text{MH}(F_{\text{leap}} \circ R) \circ \text{Refresh}.$$

This is simplified to

$$\text{MHMC} = \text{MH}(F_{\text{leap}} \circ R) \circ \text{Refresh}$$

because $\text{Refresh} \circ R = \text{Refresh}$.

Lazy chains.

Technically, many of the convergence results actually hold for *lazy* versions of the Markov chain. Specifically, for $\ell \in [0, 1]$, the ℓ -**lazy** version of a Markov chain replaces its transition kernel P with the modified kernel P_ℓ given by

$$P_\ell(x, dy) = (1 - \ell) P(x, dy) + \ell \delta_x(dy).$$

The laziness condition is familiar from the study of discrete-time Markov chains on discrete state spaces, in which laziness is useful for avoiding periodic behavior. For the remainder of this chapter, we will generally be considering $\frac{1}{2}$ -lazy versions of the Metropolis–Hastings chains without explicitly mentioning this. In any case, this modification only multiplies the mixing time by a factor of 2, so it does not significantly alter the results.

Cold start vs. warm start.

When discussing Metropolis–Hastings algorithms, we must distinguish between convergence rates when initialized at a **cold start**, vs. a **warm start**. These terms are not precisely defined, but loosely speaking a cold start refers to an easily computable distribution which works well uniformly over the class of target distributions under consideration. In this section, a cold start usually refers to the $\text{normal}(0, \beta^{-1} I_d)$ distribution, where β is the smoothness of V and we assume that the minimizer of V is 0. On the other hand, a *warm start* is a distribution which is already somewhat close to the target π ; for this section, it can be taken to mean a distribution μ_0 such that $\chi^2(\mu_0 \parallel \pi) = O(1)$. Unsurprisingly, the rates are faster with a warm start.

One obvious benefit of a warm start is simply that we start closer to stationarity, so it takes less time to approach π . This aspect of the situation is similar to the discussion in Section 1.5. Basically, the simplest way to study Markov chains is via spectral theory, which is related to Poincaré inequalities and hence to the chi-squared divergence at initialization. Since the initial chi-squared divergence tends to be quite large—of order $\exp(d)$ —the resulting analysis incurs slow mixing. This source of slow mixing, however, is sometimes illusory, as it can be addressed via stronger methods such as a log-Sobolev inequality. For example, in the case of LMC, the guarantee of Theorem 4.2.7 only improves by a logarithmic factor if we assume the warm start condition $\text{KL}(\mu_0 \parallel \pi) = O(1)$.

On the other hand, for Metropolis–Hastings algorithms, we will see another crucial benefit of a warm start: *it permits us choose a larger step size while maintaining a high acceptance probability*. Indeed, let us perform a thought experiment: for each step size h , let $A_h \subseteq \mathbb{R}^d$ denote the region of

space where a proposal starting from A_h has high acceptance probability. If the choice of step size h_{cold} allows the algorithm to succeed from a cold start, it means that the algorithm should begin accepting the proposals from a potentially poor initialization, so corresponding acceptance region is nearly full: $A_{h_{\text{cold}}} \approx \mathbb{R}^d$. However, this condition on the step size may be too stringent, and optimistically we could hope that it suffices to take a step size h_{warm} for which the corresponding acceptance region has high probability at stationarity: $\pi(A_{h_{\text{warm}}}) \approx 1$. In this case, in order for the algorithm to be initialized in $A_{h_{\text{warm}}}$, we require $A_{h_{\text{warm}}}$ to also have high probability under μ_0 : $\mu_0(A_{h_{\text{warm}}}) \approx 1$. In turn, this is implied by the warm start condition $\chi^2(\mu_0 \parallel \pi) = O(1)$. It turns out that the step size thresholds h_{cold} and h_{warm} can differ substantially, and as we shall see, $h_{\text{cold}} \asymp 1/d$ and $h_{\text{warm}} \asymp 1/d^{1/2}$ for MALA. This latter effect is more fundamental and necessitates carefully delineating the two regimes (cold vs. warm).

State-of-the-art results.

We now give the current state-of-the-art convergence guarantees for the Metropolis–Hastings algorithms that we have introduced.

Theorem 7.3.4 (Cold start case, [Dwivedi et al. \(2019\)](#); [Chen et al. \(2020\)](#); [Lee et al. \(2020\)](#); [Andrieu et al. \(2024\)](#)). Suppose that the target $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and $\nabla V(0) = 0$. Consider the following Metropolis–Hastings algorithms initialized at $\text{normal}(0, \beta^{-1} I_d)$ and with an appropriately tuned choice of parameters.

1 MRW outputs a measure μ_N satisfying $\sqrt{\chi^2(\mu_N \parallel \pi)} \leq \varepsilon$ after

$$N = \tilde{O}\left(\kappa d \text{polylog} \frac{1}{\varepsilon}\right) \quad \text{iterations}.$$

2 MALA outputs a measure μ_N satisfying $\sqrt{\chi^2(\mu_N \parallel \pi)} \leq \varepsilon$ after

$$N = \tilde{O}\left(\kappa d \text{polylog} \frac{1}{\varepsilon}\right) \quad \text{iterations}.$$

If we focus on the dimension dependence, all of the results incur at least a linear dependence on d . Next, we present the results under a warm start.

Theorem 7.3.5 (Warm start case, [Chewi et al. \(2021\)](#); [Wu et al. \(2021\)](#); [Chen and Gtarmiry \(2023a,b\)](#)). Suppose that $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$. Consider the following Metropolis–Hastings algorithms initialized at a distribution satisfying $\chi^2(\mu_0 \parallel \pi) = O(1)$ and with an appropriately tuned choice of parameters.

1 MALA outputs a measure μ_N satisfying $\|\mu_N - \pi\|_{\text{TV}} \leq \varepsilon$ after

$$N = \tilde{O}\left(\kappa d^{1/2} \text{polylog} \frac{1}{\varepsilon}\right) \quad \text{iterations}.$$

2 Assume in addition that $\nabla^2 V$ is Lipschitz in the Hilbert–Schmidt norm: $\|\nabla^2 V(x) - \nabla^2 V(y)\|_{\text{HS}} \leq \zeta \|x - y\|$ for all $x, y \in \mathbb{R}^d$. Define the scale-invariant quantity $\bar{\kappa} := \zeta / \alpha^{3/2}$. Then, MHMC outputs a measure μ_N satisfying $\|\mu_N - \pi\|_{\text{TV}} \leq \varepsilon$ after

$$N = \tilde{O}\left((\kappa + \bar{\kappa}^{2/3}) d^{1/4} \text{polylog} \frac{1}{\varepsilon}\right) \quad \text{iterations}.$$

In keeping with the theme of the rest of the book, we will focus on the results which do not assume higher-order smoothness. As we shall see, the results for MALA in both the cold and warm start cases are sharp in a suitable sense. In other words, the improved guarantees under a warm start are not a failure of the analysis technique, but are intrinsic to the problem. Hence, it is also important to study how we can *obtain* a warm start, i.e., a measure satisfying $\chi^2(\mu_0 \parallel \pi) = O(1)$. In Chapter 6, or more specifically [Corollary 6.3.2](#), we accomplished this goal via our analysis of [ULMC](#) in Rényi divergence. Critically, the cost of computing the warm start is not larger than the cost of running MALA afterwards using the warm start, leading to the following high-accuracy guarantee.

Corollary 7.3.6. *Suppose that $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and $\nabla V(0) = 0$. There is a sampler which outputs μ satisfying $\sqrt{\text{KL}(\mu \parallel \pi)} \leq \varepsilon$ using*

$$N = \tilde{O}\left(\kappa d^{1/2} \text{polylog} \frac{1}{\varepsilon}\right) \quad \text{first-order queries}.$$

This result will be extended through the use of the proximal sampler in Chapter 8: we will relax the assumption of strong log-concavity to a log-Sobolev inequality; we will treat the case of composite sampling; and we will show how to obtain this result using the proximal sampler alone.

Recently, it was shown that a warm start can be computed for MHMC as well.

Theorem 7.3.7 ([Zhang et al. \(2026\)](#)). *Suppose that $\pi \propto \exp(-V)$ satisfies $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$, $\nabla^2 V$ is ζ -Lipschitz in the Hilbert–Schmidt norm, and $\nabla V(0) = 0$. Then, there is an algorithm which outputs μ_{warm} satisfying $\chi^2(\mu_{\text{warm}} \parallel \pi) \leq O(1)$ using*

$$N = \tilde{O}\left((\kappa^{3/2} + \kappa^{1/2} \bar{\kappa}^{1/2}) d^{1/4}\right) \quad \text{first-order queries}.$$

Therefore, there is a sampler which outputs μ satisfying $\chi^2(\mu \parallel \pi) \leq \varepsilon^2$ using

$$N = \tilde{O}\left((\kappa^{3/2} + \kappa^{1/2} \bar{\kappa}^{1/2}) d^{1/4} + (\kappa + \bar{\kappa}^{2/3}) d^{1/4} \text{polylog} \frac{1}{\varepsilon}\right) \quad \text{first-order queries}.$$

This result justifies that MHMC achieves an improved dimensional scaling of $d^{1/4}$ as compared with MALA, albeit under stronger assumptions (indeed, in many settings, ζ can scale with the dimension as well). The results for MHMC lie outside the scope of this book, and the goal for the rest of the chapter is to prove the convergence results for MALA.

7.4 Markov Chains in Discrete Time

As discussed in the introduction to this chapter, the key advantage of Metropolis–Hastings algorithms is that they are *unbiased* and hence lead to high-accuracy algorithms. In order to prove complexity bounds that scale as $\text{polylog}(1/\varepsilon)$, where ε is the target accuracy, it is important that we do not simply bound the distance between MALA and, e.g., the continuous-time Langevin diffusion, as we did in Chapter 4. This is not to say that tools from Chapter 4 are completely irrelevant, only that we must first develop some new techniques for studying discrete-time Markov chains.

7.4.1 Markov Semigroup Theory

Let P be a Markov kernel. It generates a discrete-time semigroup $(P^n)_{n \in \mathbb{N}}$, and some of the ideas from Markov semigroup theory (Section 1.2) can be adapted to the present context.

Generator.

We define the **generator** of the semigroup to be the operator $\mathcal{L} := P - \text{id}$, acting on $L^2(\pi)$ via $Pf(x) := \int f(y) P(x, dy)$ (where π is the stationary distribution for P). Note that since the operator norm of P is at most 1, then $P - \text{id}$ is always a negative operator (similarly to the infinitesimal generator $\mathcal{L}$ from Section 1.2).

Reversibility.

We defined reversibility in Section 7.2 and showed that Metropolis–Hastings algorithms are reversible w.r.t. the target distribution π . For the rest of the section, we will focus on reversible Markov chains.

Spectral gap.

The **spectral gap** of P is the largest $\lambda > 0$ such that for all $f \in L^2(\pi)$ with $\mathbb{E}_\pi f = 0$,

$$\langle f, (-\mathcal{L})f \rangle_{L^2(\pi)} \geq \lambda \|f\|_{L^2(\pi)}^2.$$

Equivalently, if (X_0, X_1) are two successive iterates of the chain started at stationarity, it is equivalent to require

$$2\lambda \text{var } f(X_0) \leq \mathbb{E}[|f(X_1) - f(X_0)|^2]. \quad (7.4.1)$$

In analogy with Section 1.2, we also say that P satisfies a **Poincaré inequality** with constant $1/\lambda$. We already saw in Section 2.6 that a Poincaré inequality is implied by a lower bound on the coarse Ricci curvature.

The right-hand side of (7.4.1) can be interpreted as a **Dirichlet energy**,

$$\mathcal{E}(f, f) := \langle f, (-\mathcal{L})f \rangle_{L^2(\pi)},$$

and the Markov chain can be viewed as an $L^2(\pi)$ gradient descent on the Dirichlet energy; see [Exercise 7.3](#). We have the following convergence result.

Theorem 7.4.2 (Mixing under a spectral gap). *Suppose that the spectral gap of P is $\lambda > 0$. Then, for the law $(\mu_n)_{n \in \mathbb{N}}$ of the iterates of the $\frac{1}{2}$ -lazy version of P , we have*

$$\chi^2(\mu_N \parallel \pi) \leq \exp(-\lambda N) \chi^2(\mu_0 \parallel \pi).$$

Modified log-Sobolev inequality.

We say that P satisfies a **modified log-Sobolev inequality (MLSI)** with constant C_{MLSI} if for all $f \in L^2(\pi)$ with $f \geq 0$,

$$\text{ent}_\pi f \leq \frac{C_{\text{MLSI}}}{2} \mathcal{E}(f, \log f).$$

We have already encountered this inequality as [Definition 1.2.25](#), although there we simply called it the log-Sobolev inequality. In the context of discrete Markov processes, however, since the chain rule

fails and the different variants of the log-Sobolev inequality are no longer equivalent, it is worth being careful about the terminology.

It is trickier to deduce entropy decay from the MLSI in discrete time, and to avoid this issue we shall work in continuous time instead. The Markov kernel P gives rise to the generator $\mathcal{L} := P - \text{id}$, which in turn generates a *continuous-time* semigroup $(P_t)_{t \geq 0}$ via $P_t := \exp(t\mathcal{L})$. Note that the generator of $(P_t)_{t \geq 0}$ is also $\mathcal{L}$ and hence the Dirichlet energy for $(P_t)_{t \geq 0}$ coincides with the Dirichlet energy for $(P^n)_{n \in \mathbb{N}}$. Now, if we apply the calculation (1.2.24) to the semigroup $(P_t)_{t \geq 0}$, we obtain exponential decay of the KL divergence under an MLSI:

$$\text{KL}(\mu P_t \parallel \pi) \leq \exp\left(-\frac{2t}{C_{\text{MLSI}}}\right) \text{KL}(\mu \parallel \pi),$$

see [Theorem 1.2.26](#).

Moreover, the continuous-time semigroup $(P_t)_{t \geq 0}$ can be simulated. Namely, let $(\tau_n)_{n \in \mathbb{N}^+} \stackrel{\text{i.i.d.}}{\sim} \text{exponential}(1)$, $T_n := \sum_{k=1}^n \tau_k$, and consider the following algorithm. Initialize at $X_0 \sim \mu_0$, and for $n = 0, 1, 2, \dots$, let $X_{T_{n+1}} \sim P(X_{T_n}, \cdot)$, so that $(X_{T_n})_{n \in \mathbb{N}}$ are the iterates of the discrete-time Markov chain with kernel P . Also, for $T_n \leq t < T_{n+1}$, set $X_t := X_{T_n}$. This yields a continuous-time Markov process $(X_t)_{t \geq 0}$, and one can check that the associated Markov semigroup is exactly $(P_t)_{t \geq 0}$. Moreover, by concentration of i.i.d. sums, it holds that $T_n \approx n$, so that if the semigroup $(P_t)_{t \geq 0}$ requires time T_{mix} in order to mix to a desired level of accuracy, then the algorithm which simulates $(X_t)_{t \geq 0}$ requires $\approx T_{\text{mix}}$ iterations to reach the same level of mixing.

This argument can even be made rigorous, using concentration inequalities for the Poisson random variable, in order to argue that a MLSI for P implies a mixing time bound (in total variation distance, say) for the discrete-time chain $(P^n)_{n \in \mathbb{N}}$. We omit the details and content ourselves with the knowledge that an MLSI for P at least leads to the existence of an implementable algorithm (simulating $(X_t)_{t \geq 0}$) with good mixing.

We leave the converse implication (that entropy decay for the discrete-time Markov chain generated by P implies an MLSI for P) as [Exercise 7.4](#).

7.4.2 Conductance

Unfortunately, it is usually quite challenging to prove either a Poincaré inequality or a modified log-Sobolev inequality for discrete-time Markov chains, which motivates the use of conductance.

The **conductance** of P is the greatest number $\mathfrak{c} > 0$ such that for all events $A \subseteq \mathbb{R}^d$,

$$\int_A P(x, A^c) \pi(dx) \geq \mathfrak{c} \pi(A) \pi(A^c).$$

A small conductance implies the presence of *bottlenecks* in the space: subsets A of the state space from which it is difficult for the Markov chain to exit. On the other hand, it is a remarkable fact that once the presence of these bottlenecks is eliminated, then there is a positive spectral gap. This is the content of a celebrated result of Cheeger.

Theorem 7.4.3 (Cheeger's inequality, [Lawler and Sokal \(1988\)](#)). *The conductance $\mathfrak{c}$ and the spectral gap λ satisfy the inequalities*

$$\frac{1}{8} \mathfrak{c}^2 \leq \lambda \leq \mathfrak{c}.$$

Both inequalities are sharp up to constants. The upper bound on λ is immediate (see [Exercise 7.3](#)), so we focus on the lower bound. We begin by reformulating the conductance as a functional inequality.

Lemma 7.4.4 (Functional form of conductance). *Let the conductance of the chain be $\mathfrak{c} > 0$. Then, for all $f \in L^1(\pi)$,*

$$\mathbb{E}_\pi |f - \text{med}_\pi f| \leq \frac{1}{\mathfrak{c}} \mathbb{E} |f(X_1) - f(X_0)|, \quad (7.4.5)$$

where (X_0, X_1) are two successive iterates of the chain started at stationarity.

Proof Let $m := \text{med}_\pi f$. Then, by reversibility,

$$\begin{aligned} \mathbb{E} |f(X_1) - f(X_0)| &= \iint |f(x_1) - f(x_0)| P(x_0, dx_1) \pi(dx_0) \\ &= 2 \iint \mathbb{1}\{f(x_1) > f(x_0)\} [f(x_1) - f(x_0)] P(x_0, dx_1) \pi(dx_0) \\ &= 2 \iiint \mathbb{1}\{f(x_1) > t \geq f(x_0)\} P(x_0, dx_1) \pi(dx_0) dt \\ &= 2 \int \left(\int_{\{f \leq t\}^c} P(x_0, \{f \leq t\}^c) \pi(dx_0) \right) dt \\ &\geq 2\mathfrak{c} \int \pi(\{f \leq t\}) \pi(\{f > t\}) dt \\ &\geq \mathfrak{c} \left(\int_{-\infty}^m \pi(\{f \leq t\}) dt + \int_m^\infty \pi(\{f > t\}) dt \right) = \mathfrak{c} \mathbb{E}_\pi |f - m|. \quad \square \end{aligned}$$

Compare this with the relationship between the Cheeger isoperimetric inequality and the L^1 – L^1 Poincaré inequality in [Theorem 2.4.18](#). Indeed, the trick above of passing to the level sets of f is the discrete version of the coarea inequality ([Theorem 2.4.16](#)).

Recall also that an L^1 – L^1 Poincaré inequality implies an L^2 – L^2 Poincaré inequality with $C_{2,2} \leq C_{1,1}$, see [Proposition 2.4.21](#). On the other hand, $C_{2,2}$ is the square root of the usual Poincaré constant, $C_{2,2} = 1/\sqrt{\lambda}$, where λ is the spectral gap. To prove Cheeger's inequality, we are going to follow the same principle in discrete time. This is exactly the source of the square in the lower bound $\lambda \gtrsim \mathfrak{c}^2$ of Cheeger's inequality.

Proof of Cheeger's inequality ([Theorem 7.4.3](#)) For $f \in L^2(\pi)$, let $m := \text{med}_\pi f$. Then, 0 is a median of $(f - m) |f - m|$, so by [Lemma 7.4.4](#), the elementary inequality $|a| |a| - b |b| \leq |a - b| (|a| + |b|)$ for $a, b \in \mathbb{R}$, and Cauchy–Schwarz,

$$\begin{aligned} \mathfrak{c} \mathbb{E}_\pi [|f - m|^2] &\leq \mathbb{E} |(f(X_1) - m) |f(X_1) - m| - (f(X_0) - m) |f(X_0) - m|| \\ &\leq \mathbb{E} [|f(X_1) - f(X_0)| (|f(X_0) - m| + |f(X_1) - m|)] \\ &\leq 2\sqrt{\mathbb{E} [|f(X_1) - f(X_0)|^2] \mathbb{E} [|f(X_0) - m|^2]} = \sqrt{8 \mathcal{E}(f, f) \mathbb{E}_\pi [|f - m|^2]}, \end{aligned}$$

which yields

$$\text{var}_\pi f \leq \mathbb{E}_\pi [|f - m|^2] \leq \frac{8 \mathcal{E}(f, f)}{\mathfrak{c}^2}. \quad \square$$

A lower bound on the conductance via overlaps.

At this stage, it may not seem that we have gained anything by moving from the spectral gap to the conductance. We now introduce a key lemma, which provides a tractable lower bound on the conductance in terms of two geometric quantities: a Cheeger isoperimetric inequality for target π (we introduced this inequality in Section 2.4.3), and overlap bounds on the Markov chain. Recall that an α -strongly log-concave measure π satisfies the Cheeger isoperimetric inequality with $\text{Ch} \lesssim 1/\sqrt{\alpha}$ (Corollary 2.4.23).

Lemma 7.4.6 (Overlap lemma). *Assume the following:*

- 1 *The target π satisfies a Cheeger isoperimetric inequality with constant $\text{Ch} > 0$.*
 - 2 *There exists $r > 0$ such that for any points $x, y \in \mathbb{R}^d$ with $\|x - y\| \leq r$, it holds that $\|P(x, \cdot) - P(y, \cdot)\|_{\text{TV}} \leq \frac{1}{2}$.*
- Then, $\mathfrak{c} \geq \min\{\frac{1}{8}, \frac{r}{32\text{Ch}}\}$.*

Proof Let $A_0 \subseteq \mathbb{R}^d$; for symmetry of notation, write $A_1 := A_0^c$. By reversibility,

$$\int_{A_0} P(x, A_1) \pi(\mathrm{d}x) = \int_{A_1} P(y, A_0) \pi(\mathrm{d}y) = \frac{1}{2} \left(\int_{A_0} P(x, A_1) \pi(\mathrm{d}x) + \int_{A_1} P(y, A_0) \pi(\mathrm{d}y) \right).$$

We want to lower bound this by a constant times $\pi(A_0) \pi(A_1)$.

Define bad sets and a good set:

$$\begin{aligned} B_0 &:= \{x \in A_0 \mid P(x, A_1) < \frac{1}{4}\}, \\ B_1 &:= \{y \in A_1 \mid P(y, A_0) < \frac{1}{4}\}, \\ G &:= \mathbb{R}^d \setminus (B_0 \cup B_1). \end{aligned}$$

We can assume that $\pi(B_0) \geq \pi(A_0)/2$ and $\pi(B_1) \geq \pi(A_1)/2$. Indeed, if we have, e.g., $\pi(B_0) \leq \pi(A_0)/2$, then $\pi(A_0 \setminus B_0) \geq \pi(A_0)/2$, and

$$\int_{A_0} P(x, A_1) \pi(\mathrm{d}x) \geq \int_{A_0 \setminus B_0} P(x, A_1) \pi(\mathrm{d}x) \geq \frac{1}{4} \pi(A_0 \setminus B_0) \geq \frac{1}{8} \pi(A_0).$$

Next, suppose that $x \in B_0$ and $y \in B_1$. Then, $P(x, A_0) \geq \frac{3}{4}$, whereas $P(y, A_0) < \frac{1}{4}$. It follows that $\|P(x, \cdot) - P(y, \cdot)\|_{\text{TV}} > \frac{1}{2}$. By our second assumption, $\|x - y\| > r$. This shows that $B_1 \subseteq (B_0^r)^c$, or $B_1^c \supseteq B_0^r$. On the other hand, $G = B_0^c \cap B_1^c \supseteq B_0^r \setminus B_0$. Integrating the isoperimetric inequality in (2.4.14) shows that

$$\pi(G) \geq \pi(B_0^r) - \pi(B_0) \geq \frac{r}{\text{Ch}} \pi(B_0) \pi((B_0^r)^c) \geq \frac{r}{\text{Ch}} \pi(B_0) \pi(B_1) \geq \frac{r}{4\text{Ch}} \pi(A_0) \pi(A_1).$$

Hence,

$$\begin{aligned}
& \frac{1}{2} \left(\int_{A_0} P(x, A_1) \pi(dx) + \int_{A_1} P(y, A_0) \pi(dy) \right) \\
& \geq \frac{1}{2} \left(\int_{A_0 \cap G} P(x, A_1) \pi(dx) + \int_{A_1 \cap G} P(y, A_0) \pi(dy) \right) \\
& \geq \frac{1}{8} (\pi(A_0 \cap G) + \pi(A_1 \cap G)) = \frac{1}{8} \pi(G) \geq \frac{r}{32 \text{Ch}} \pi(A_0) \pi(A_1). \quad \square
\end{aligned}$$

From conductance to s -conductance.

Unfortunately, the framework that we have developed so far is not flexible enough to study MALA. In particular, requiring that the second condition in [Lemma 7.4.6](#) hold for *all* pairs of points $x, y \in \mathbb{R}^d$ is rather restrictive, especially because there are many points which we are unlikely to ever visit in the course of running the sampling algorithm. To address these issues, many variants of conductance have been proposed in the literature. Here we will introduce only one other variant, the *s -conductance*, which seems reasonably flexible.

For $s \in [0, \frac{1}{2}]$, the *s -conductance* of P is the largest $\mathfrak{c}_s > 0$ such that for all events $A \subseteq \mathbb{R}^d$,

$$\int_A P(x, A^c) \pi(dx) \geq \mathfrak{c}_s (\pi(A) - s) (\pi(A^c) - s).$$

Observe that if $\pi(A) \leq s$, then the above inequality holds trivially. Hence, this definition allows us to restrict attention to events which are reasonably probable under π .

For the conductance, we had Cheeger's inequality which relates conductance to the spectral gap and ultimately to convergence. For the s -conductance, the following theorem is an appropriate substitute.

Theorem 7.4.7 ([Lovász and Simonovits \(1993, Corollary 1.6\)](#)). *For any $0 < s \leq \frac{1}{2}$, let*

$$\Delta_s := \sup\{|\mu_0(A) - \pi(A)| : A \subseteq \mathbb{R}^d, \pi(A) \leq s\}.$$

Then, the law μ_N of the N -th iterate of a Markov chain with s -conductance $\mathfrak{c}_s$ and initialized at μ_0 satisfies

$$\|\mu_N - \pi\|_{\text{TV}} \leq \Delta_s + \frac{\Delta_s}{s} \exp\left(-\frac{\mathfrak{c}_s^2 N}{2}\right).$$

In particular,

$$\|\mu_N - \pi\|_{\text{TV}} \leq \sqrt{s \chi^2(\mu_0 \parallel \pi)} + \sqrt{\frac{\chi^2(\mu_0 \parallel \pi)}{s}} \exp\left(-\frac{\mathfrak{c}_s^2 N}{2}\right).$$

Proof The first statement is from [Lovász and Simonovits \(1993, Corollary 1.6\)](#) and the proof is not particularly straightforward. We sketch an alternative approach in [Exercise 7.5](#).

The second statement follows from the first: indeed, for $A \subseteq \mathbb{R}^d$ with $\pi(A) \leq s$,

$$|\mu_0(A) - \pi(A)| = \left| \int \mathbb{1}_A d(\mu_0 - \pi) \right| = \left| \int \mathbb{1}_A \left(\frac{d\mu_0}{d\pi} - 1 \right) d\pi \right| \leq \sqrt{\pi(A) \chi^2(\mu_0 \parallel \pi)}$$

so that $\Delta_s \leq \sqrt{s \chi^2(\mu_0 \parallel \pi)}$. □

This result says that if $s = \varepsilon^2 / 4\chi^2(\mu_0 \parallel \pi)$, then we obtain $\|\mu_N - \pi\|_{\text{TV}} \leq \varepsilon$ after

$$N = O\left(\frac{1}{\varepsilon_s^2} \log \frac{\chi^2(\mu_0 \parallel \pi)}{\varepsilon^2}\right) \quad \text{iterations}.$$

Although this mixing time guarantee is stated in the TV metric, the guarantee can often be “boosted” to stronger metrics, see [Exercise 7.7](#).

The key advantage of the s -conductance is that it allows for a version of the key lemma with weaker assumptions; the proof is left as [Exercise 7.6](#).

Lemma 7.4.8 (Overlap lemma for s -conductance). *Assume the following:*

- 1 The target π satisfies a Cheeger isoperimetric inequality with constant $\text{Ch} > 0$.
- 2 There exists $r \in [0, \text{Ch}]$ and an event $E \subseteq \mathbb{R}^d$ with probability $\pi(E) \geq 1 - \frac{rs}{16\text{Ch}}$ such that

$$\forall x, y \in E, \|x - y\| \leq r \implies \|P(x, \cdot) - P(y, \cdot)\|_{\text{TV}} \leq \frac{1}{2}.$$

Then, $\varepsilon_s \gtrsim r/\text{Ch}$.

7.5 Analysis of MALA for a Cold Start

Using the tools we have developed, we now proceed to analyze the mixing time of MALA under the assumptions of [Theorem 7.3.4](#). However, we will not prove the full strength of the result in [Theorem 7.3.4](#); at the end of this section, we will indicate the extra steps needed to reach [Theorem 7.3.4](#).

Basic decomposition.

The overall plan is to lower bound the s -conductance using the key lemma ([Lemma 7.4.8](#)), which then upper bounds the mixing time via [Theorem 7.4.7](#). By strong log-concavity of π , the first hypothesis of [Lemma 7.4.8](#) is verified, so it remains to bound the overlaps. For a kernel P , we use the shorthand $P_x := P(x, \cdot)$.

By the triangle inequality, we have the decomposition

$$\|P_x - P_y\|_{\text{TV}} \leq \|P_x - Q_x\|_{\text{TV}} + \|Q_x - Q_y\|_{\text{TV}} + \|P_y - Q_y\|_{\text{TV}}. \quad (7.5.1)$$

The middle term $\|Q_x - Q_y\|_{\text{TV}}$ measures the overlap for the proposal kernel, and we will shortly see that this term is easy to bound. Then, controlling the first and third terms essentially amounts to lower bounding the acceptance probability of MALA, since the only difference between P and Q is the Metropolis–Hastings filter.

Overlap of the proposal kernel.

Lemma 7.5.2. *For $x \in \mathbb{R}^d$, let $Q_x := \text{normal}(x - h \nabla V(x), 2h I_d)$, and assume that $\|\nabla^2 V\|_{\text{op}} \leq \beta$. Then, provided $h \leq 1/\beta$, we have*

$$\|Q_x - Q_y\|_{\text{TV}} \leq \frac{\|x - y\|}{\sqrt{2h}}.$$

Proof By Pinsker's inequality,

$$\|Q_x - Q_y\|_{\text{TV}}^2 \leq \frac{1}{2} \text{KL}(Q_x \| Q_y) = \frac{\|x - h \nabla V(x) - y + h \nabla V(y)\|^2}{8h} \leq \frac{\|x - y\|^2}{2h}$$

where the last inequality uses the fact that $\text{id} - h \nabla V$ is 2-Lipschitz. $\square$

Control of the acceptance probability.

Next, consider the term $\|P_x - Q_x\|_{\text{TV}}$. Computing this term is slightly tricky because P_x has an atom at x , but in the end we obtain

$$\begin{aligned} \|P_x - Q_x\|_{\text{TV}} &= \frac{1}{2} \left[\underbrace{1 - \int Q(x, dy) A(x, y)}_{\text{from the atom of } P_x} + \int_{\mathbb{R}^d \setminus \{x\}} |P(x, y) - Q(x, y)| dy \right] \\ &= \frac{1}{2} \left[1 - \int Q(x, dy) A(x, y) + \int Q(x, dy) \{1 - A(x, y)\} dy \right] \\ &= 1 - \int Q(x, dy) A(x, y). \end{aligned} \quad (7.5.3)$$

This has a very clear interpretation: it is the probability that the proposed move starting at x is rejected. If we let $\xi \sim \text{normal}(0, I_d)$ and $Y := x - h \nabla V(x) + \sqrt{2h} \xi$, we want a lower bound on the quantity $\mathbb{E} A(x, Y)$, which comes from Markov's inequality:

$$\mathbb{E} A(x, Y) = \mathbb{E} \min\left\{1, \frac{\pi(Y) Q(Y, x)}{\pi(x) Q(x, Y)}\right\} \geq \lambda \mathbb{P}\left\{\frac{\pi(Y) Q(Y, x)}{\pi(x) Q(x, Y)} \geq \lambda\right\} \quad \text{for all } 0 < \lambda < 1.$$

The approach now is to write out the ratio more explicitly, and then carefully group together and bound the terms.

Explicitly, we have

$$\frac{\pi(Y) Q(Y, x)}{\pi(x) Q(x, Y)} = \exp\left(-V(Y) - \frac{\|x - Y + h \nabla V(Y)\|^2}{4h} + V(x) + \frac{\|Y - x + h \nabla V(x)\|^2}{4h}\right).$$

After some careful algebra,

$$4 \log \frac{\pi(Y) Q(Y, x)}{\pi(x) Q(x, Y)} = h \{\|\nabla V(x)\|^2 - \|\nabla V(Y)\|^2\} \quad (7.5.4)$$

$$- 2 \{V(Y) - V(x) - \langle \nabla V(x), Y - x \rangle\} \quad (7.5.5)$$

$$+ 2 \{V(x) - V(Y) - \langle \nabla V(Y), x - Y \rangle\}. \quad (7.5.6)$$

The terms are grouped to more easily apply the strong convexity and smoothness of V . It yields

$$(7.5.5) \geq -\beta \|x - Y\|^2 \quad \text{and} \quad (7.5.6) \geq \alpha \|x - Y\|^2 \geq 0.$$

Also, for $h \leq 1/\beta$,

$$\begin{aligned} (7.5.4) &= h \langle \nabla V(x) - \nabla V(Y), \nabla V(x) + \nabla V(Y) \rangle \\ &\geq -h \|\nabla V(x) - \nabla V(Y)\| \|\nabla V(x) + \nabla V(Y)\| \\ &\geq -\beta h \|x - Y\| (2 \|\nabla V(x)\| + \beta \|x - Y\|) \geq -\beta h^2 \|\nabla V(x)\|^2 - 2\beta \|x - Y\|^2. \end{aligned}$$

Therefore,

$$\log \frac{\pi(Y) Q(Y, x)}{\pi(x) Q(x, Y)} \gtrsim -\beta h^2 \|\nabla V(x)\|^2 - \beta \|x - Y\|^2 \gtrsim -\beta h^2 \|\nabla V(x)\|^2 - \beta h \|\xi\|^2.$$

At this stage, observe that we cannot lower bound this quantity (with high probability) *uniformly* over x , since $\|\nabla V(x)\| \rightarrow \infty$ as $\|x\| \rightarrow \infty$. This is why it is helpful to restrict to x belonging to some high-probability event E , which is ultimately achieved by working with s -conductance rather than conductance.

By standard concentration bounds, $\|\xi\|^2 \leq 2d$ with probability at least $1 - \exp(-d/2)$. Also, let $E_\delta := \{\|\nabla V\| \leq \sqrt{\beta(d + \log(1/\delta))}\}$. It follows that for all $x \in E_\delta$, if we take $h \lesssim 1/(\beta d + \beta \sqrt{\log(1/\delta)})$ with a sufficiently small constant, then

$$\mathbb{E} A(x, Y) \geq \frac{11}{12} \mathbb{P}\left\{\frac{\pi(Y) Q(Y, x)}{\pi(x) Q(x, Y)} \geq \frac{11}{12}\right\} \geq \frac{11}{12} \left(1 - \exp(-\frac{d}{2})\right) \geq \frac{5}{6},$$

for sufficiently large d . Hence, for $x \in E_\delta$, we have $\|P_x - Q_x\|_{\text{TV}} \leq \frac{1}{6}$.

Completing the analysis.

We have shown: if the step size satisfies $h \lesssim 1/(\beta d + \beta \sqrt{\log(1/\delta)})$, then for all $x, y \in E_\delta$ with $\|x - y\| \leq r$,

$$\|P_x - P_y\|_{\text{TV}} \leq \|P_x - Q_x\|_{\text{TV}} + \|Q_x - Q_y\|_{\text{TV}} + \|P_y - Q_y\|_{\text{TV}} \leq \frac{1}{6} + \frac{r}{\sqrt{2}h} + \frac{1}{6} \leq \frac{1}{2}$$

provided we take $r = \sqrt{2}h/6$. Applying [Lemma 7.4.8](#) (assuming $h \lesssim 1/\alpha$), we deduce that $c_s \gtrsim \sqrt{\alpha h}$ provided $\pi(E_\delta) \geq 1 - c_0 s \sqrt{\alpha h}$, where $c_0 > 0$ is a universal constant. Since we want the step size h to be as large as possible, we take $h \asymp 1/(\beta d + \beta \sqrt{\log(1/\delta)})$, where δ is chosen to satisfy $\pi(E_\delta^c) \leq s/(\kappa d + \kappa \log^{1/2}(1/\delta))^{1/2}$ and $s \asymp \varepsilon^2/\chi^2(\mu_0 \parallel \pi)$. By the gradient concentration inequality in [Lemma 6.2.7](#), $\pi(E_\delta^c) \leq \delta$, so we can choose $B = \tilde{O}(\log(\chi^2(\mu_0/\pi)/\varepsilon^2))$. The final mixing time bound implied by [Theorem 7.4.7](#) is then $\tilde{O}(\kappa d \log^2(\chi^2(\mu_0 \parallel \pi)/\varepsilon^2))$.

From warm start to cold start.

The factor of $\log \chi^2(\mu_0 \parallel \pi)$ in the bound incurs additional dimension dependence under a cold start, since with a Gaussian initialization we can only show $\chi^2(\mu_0 \parallel \pi) \leq \kappa^{O(d)}$. The problem is that the conductance-based analysis relies upon Poincaré-type inequalities, instead of log-Sobolev inequalities. To address this issue, we can replace the assumption of a Cheeger isoperimetric inequality with a *Gaussian isoperimetric inequality* (see [Section 2.4.5](#)). The essential difference is that under a Cheeger isoperimetric inequality, as $p := \pi(A) \searrow 0$ we have $\pi^+(A) \gtrsim p$, whereas under a Gaussian isoperimetric inequality we have $\pi^+(A) \gtrsim p \sqrt{\log(1/p)}$. Using this stronger assumption, [Chen et al. \(2020\)](#) show that the dependence on the initialization can be improved to $\log \log \chi^2(\mu_0 \parallel \pi)$. A similar effect can be achieved via the *blocking conductance* ([Kannan et al., 2006](#)), which was used in [Lee et al. \(2020\)](#). We omit the details.

Lower bound.

Finally, the analysis of MALA in [Theorem 7.3.4](#) is tight, as shown in the following lower bound.

Theorem 7.5.7 (Lee et al. (2021a)). *For every choice of step size $h > 0$, there exists a target distribution $\pi \propto \exp(-V)$ on $\mathbb{R}^d$ with $I_d \leq \nabla^2 V \leq \kappa I_d$, as well as an initialization μ_0 with $\chi^2(\mu_0 \parallel \pi) \leq \exp d$, such that the number of iterations required for MALA to reach total variation at most $\frac{1}{4}$ from π is at least $\tilde{\Omega}(\kappa d)$.*

This theorem is a lower bound in the sense that all of the known proofs for MALA do not use any property of the initialization μ_0 except through $\chi^2(\mu_0 \parallel \pi)$. Thus, in order to improve the analysis of MALA under a cold start, one must use more specific properties of the initialization, or use some other modification that bypasses the lower bound (e.g., random step sizes).

7.6 Analysis of MALA for a Warm Start

We next turn towards the warm start case (Theorem 7.3.5). The improvement under a warm start was first shown in Chewi et al. (2021), with subsequent refinements in Wu et al. (2021); Chen and Garmir (2023a) via completely different techniques. In this section, we follow Chewi et al. (2021) because the proof is more conceptual and closer to the stochastic calculus arguments of Chapter 6.

We still follow the s -conductance framework of the previous section, including the basic decomposition (7.5.1). The main difference lies in the control of $\|P_x - Q_x\|_{\text{TV}}$, which was previously accomplished by lower bounding the acceptance probability. Surprisingly, the following proof never works directly with the acceptance probability, despite the fact that $\|P_x - Q_x\|_{\text{TV}}$ is precisely the rejection probability at x (see (7.5.3)).

Using the projection property.

The key insight is to use projection characterization of the Metropolis–Hastings filter (Theorem 7.2.5): the MALA kernel P is the closest kernel to the proposal Q (in an appropriate L^1 distance) among all reversible Markov chains with stationary distribution π . Concretely, for any other kernel $\bar{Q}$ which is reversible w.r.t. π ,

$$\iint_{(\mathbb{R}^d \times \mathbb{R}^d) \setminus \text{diag}} |P(x, y) - Q(x, y)| \pi(\text{d}x) \text{d}y \leq \iint_{(\mathbb{R}^d \times \mathbb{R}^d) \setminus \text{diag}} |Q(x, y) - \bar{Q}(x, y)| \pi(\text{d}x) \text{d}y.$$

Now supposing that $\bar{Q}$ has no atoms, this inequality is the same as

$$\int \|P_x - Q_x\|_{\text{TV}} \pi(\text{d}x) \leq 2 \int \|Q_x - \bar{Q}_x\|_{\text{TV}} \pi(\text{d}x). \quad (7.6.1)$$

Thus, we can indirectly bound $\|Q_x - P_x\|_{\text{TV}}$, at least on average. Moreover, there is a very natural choice of $\bar{Q}$ here: since Q is obtained from a discretization of the Langevin diffusion, we can take $\bar{Q}$ to be the continuous-time Langevin diffusion run for time h , which is indeed reversible with respect to π . The right-hand side of the above expression then simply measures the discretization error, which we have already studied in detail.

Pointwise projection property.

The projection property is not enough for our purposes, however, since it only bounds $\|P_x - Q_x\|_{\text{TV}}$ in average, whereas we really need high-probability bounds. Thankfully, we can extend the projection property to a pointwise bound.

Theorem 7.6.2 (Pointwise projection property, [Chewi et al. \(2021, Theorem 6\)](#)). *Let Q be an atomless proposal kernel and let P be the corresponding Metropolis–Hastings kernel with target π . Then, for any atomless kernel $\bar{Q}$ which is reversible with respect to π , and for every $x \in \mathbb{R}^d$,*

$$\|P_x - Q_x\|_{\text{TV}} \leq 2 \|Q_x - \bar{Q}_x\|_{\text{TV}} + \int \frac{\pi(y) \bar{Q}(y, x)}{\pi(x)} \left| \frac{Q(y, x)}{\bar{Q}(y, x)} - 1 \right| dy.$$

Consequently, for any convex increasing function $\Phi : \mathbb{R}_+ \rightarrow \mathbb{R}_+$,

$$\begin{aligned} \int \Phi(\|P_x - Q_x\|_{\text{TV}}) \pi(dx) &\leq \frac{1}{2} \int \Phi(4 \|Q_x - \bar{Q}_x\|_{\text{TV}}) \pi(dx) \\ &\quad + \frac{1}{2} \iint \Phi\left(2 \left| \frac{Q(x, y)}{\bar{Q}(x, y)} - 1 \right| \right) \bar{Q}(x, dy) \pi(dx). \end{aligned} \quad (7.6.3)$$

We will not need the inequality (7.6.3), so the proof is left as [Exercise 7.9](#). The reason why (7.6.3) is included in the theorem is because it makes it clear why we can expect the pointwise projection property to imply high-probability bounds for $\|P_x - Q_x\|_{\text{TV}}$. Note that when we integrate the projection property w.r.t. $\pi(dx)$, we recover (7.6.1) with a factor of 4 on the right-hand side instead of 2.

Proof We can write

$$\begin{aligned} \|P_x - Q_x\|_{\text{TV}} &= 1 - \int Q(x, dy) A(x, y) = \int \left[1 - \left(1 \wedge \frac{\pi(y) Q(y, x)}{\pi(x) Q(x, y)} \right) \right] Q(x, dy) \\ &\leq \int \left| 1 - \frac{\pi(y) Q(y, x)}{\pi(x) Q(x, y)} \right| Q(x, dy) \\ &\leq \int \left| 1 - \frac{\pi(y) \bar{Q}(y, x)}{\pi(x) Q(x, y)} \right| Q(x, dy) + \int \frac{\pi(y) \bar{Q}(y, x)}{\pi(x)} \left| \frac{Q(y, x)}{\bar{Q}(y, x)} - 1 \right| dy. \end{aligned}$$

Using reversibility of $\bar{Q}$, the first term is

$$\int \left| 1 - \frac{\pi(y) \bar{Q}(y, x)}{\pi(x) Q(x, y)} \right| Q(x, dy) = \int |Q(x, y) - \bar{Q}(x, y)| dy = 2 \|Q_x - \bar{Q}_x\|_{\text{TV}},$$

which completes the proof. $\square$

Applying the pointwise projection property.

Our goal is to bound $\|P_x - Q_x\|_{\text{TV}}$ for all x which lies in an event E of very high probability under π . We proceed by controlling the two terms in the pointwise projection property separately.

The first term, $\|Q_x - \bar{Q}_x\|_{\text{TV}}$, is more straightforward. It is helpful to apply Pinsker's inequality, leaving us to control $\text{KL}(\bar{Q}_x \| Q_x)$. This is precisely the kind of discretization error that we controlled via Girsanov's theorem in Section 4.4. In particular, it follows from Theorem 4.4.1 that $\|Q_x - \bar{Q}_x\|_{\text{TV}} \lesssim \beta h^{3/2} \|\nabla V(x)\| + \beta d^{1/2} h$. By the gradient concentration bound in Lemma 6.2.7, for $h \lesssim 1/\beta$, we have $\|Q_x - \bar{Q}_x\|_{\text{TV}} \lesssim \beta d^{1/2} h + \beta^{3/2} h^{3/2} \sqrt{\log(1/\delta)}$ with probability $1 - \delta$ under π . In particular, we expect a step size of $h \asymp 1/\beta d^{1/2}$ to suffice to control this term, which is encouraging.

To control the second term with high probability, it suffices to control the moments of this quantity under π : for $p \geq 1$,

$$\int \left| \int \frac{\pi(y) \bar{Q}(y, x)}{\pi(x)} \left| \frac{Q(y, x)}{\bar{Q}(y, x)} - 1 \right| dy \right|^p \pi(dx) \leq \iint \left| \frac{Q(y, x)}{\bar{Q}(y, x)} - 1 \right|^p \bar{Q}(y, dx) \pi(dy).$$

Let $\bar{\mathcal{Q}}_x$ denote the measure on path space $C([0, h]; \mathbb{R}^d)$ of the Langevin diffusion started at x . Similarly, let $\mathcal{Q}_x$ denote the same for the interpolation of LMC. By the data-processing inequality, we obtain

$$\iint \left| \frac{\mathcal{Q}(y, x)}{\bar{\mathcal{Q}}(y, x)} - 1 \right|^p \bar{\mathcal{Q}}(y, dx) \pi(dy) \leq \int \mathbb{E}^{\bar{\mathcal{Q}}_x} \left[\left| \frac{d\mathcal{Q}_x}{d\bar{\mathcal{Q}}_x} - 1 \right|^p \right] \pi(dx).$$

We can again approach this term via stochastic calculus, although the arguments become more involved. First, via Girsanov's theorem (Theorem 3.2.8),

$$\frac{d\mathcal{Q}_x}{d\bar{\mathcal{Q}}_x} = \exp \left(\underbrace{\frac{1}{\sqrt{2}} \int_0^h \langle \nabla V(x) - \nabla V(X_t), dB_t \rangle - \frac{1}{4} \int_0^h \|\nabla V(x) - \nabla V(X_t)\|^2 dt}_{:=L_h} \right),$$

where $(B_t)_{t \geq 0}$ is the $\bar{\mathcal{Q}}_x$ -Brownian motion and $(X_t)_{t \geq 0}$ is the Langevin diffusion driven by $(B_t)_{t \geq 0}$ and started at x . By applying Itô's formula (Theorem 1.1.19),

$$\exp(L_h) - 1 = \frac{1}{\sqrt{2}} \int_0^h (\exp L_t) \langle \nabla V(X_t) - \nabla V(x), dB_t \rangle.$$

We need to control the p -th moment of this stochastic integral; we can take p to be an even integer. For $p = 2$, we can apply the Itô isometry (1.1.9). For larger values of p , the appropriate substitute is the Burkholder–Davis–Gundy inequality with optimal constants.

Lemma 7.6.4 (Burkholder–Davis–Gundy, Burkholder (1973); Davis (1976)). *There is a universal constant $C > 0$ such that for any even integer $p \geq 2$,*

$$\mathbb{E} \left[\sup_{t \in [0, T]} \left| \int_0^t \langle \eta_s, dB_s \rangle \right|^p \right] \leq (Cp)^{p/2} \mathbb{E} \left[\left| \int_0^T \|\eta_t\|^2 dt \right|^{p/2} \right].$$

We have stated this inequality in a form that is convenient for our purposes, although it applies more generally to local martingales. The growth of the optimal constants with p implies that the stochastic integral has sub-Gaussian tails.

Applying this to the problem at hand,

$$\begin{aligned} \mathbb{E}^{\bar{\mathcal{Q}}_x} \left[\left| \frac{d\mathcal{Q}_x}{d\bar{\mathcal{Q}}_x} - 1 \right|^p \right] &\leq (Cp)^{p/2} \mathbb{E}^{\bar{\mathcal{Q}}_x} \left[\left| \int_0^h \exp(2L_t) \|\nabla V(X_t) - \nabla V(x)\|^2 dt \right|^{p/2} \right] \\ &\leq (C\beta^2 p)^{p/2} \mathbb{E}^{\bar{\mathcal{Q}}_x} \left[\left| \int_0^h \exp(4L_t) dt \right|^{p/4} \left| \int_0^h \|X_t - x\|^4 dt \right|^{p/4} \right] \\ &\leq (C\beta^2 p)^{p/2} \sqrt{\mathbb{E}^{\bar{\mathcal{Q}}_x} \left[\left| \int_0^h \exp(4L_t) dt \right|^{p/2} \right] \mathbb{E}^{\bar{\mathcal{Q}}_x} \left[\left| \int_0^h \|X_t - x\|^4 dt \right|^{p/2} \right]} \\ &\leq (C\beta^2 p)^{p/2} h^{p/2-1} \sqrt{\left(\mathbb{E}^{\bar{\mathcal{Q}}_x} \int_0^h \exp(2pL_t) dt \right) \left(\mathbb{E}^{\bar{\mathcal{Q}}_x} \int_0^h \|X_t - x\|^{2p} dt \right)}. \end{aligned}$$

We bound each of the remaining expectation terms in turn. We start by observing that $\mathbb{E}^{\bar{\mathcal{Q}}_x} \exp(2pL_t) = \exp((2p-1) \mathcal{R}_{2p}(\mathcal{Q}_x|_{[0,t]} | \bar{\mathcal{Q}}_x|_{[0,t]}))$, so by Theorem 6.2.1 we obtain

$$\mathbb{E}^{\bar{\mathcal{Q}}_x} \int_0^h \exp(2pL_t) dt \leq h \exp O(\beta^2 h^3 p^2 \|\nabla V(x)\|^2 + \beta^2 dh^2 p^2)$$

provided $h \lesssim 1/\beta p$. Next, by (6.2.2) and Lemma 6.2.4,

$$\begin{aligned} \mathbb{E}^{\tilde{\mathcal{Q}}_x} \int_0^h \|X_t - x\|^{2p} dt &\leq C^p h \left(h^{2p} \|\nabla V(x)\|^{2p} + \mathbb{E}^{\tilde{\mathcal{Q}}_x} \sup_{t \in [0, h]} \|B_t\|^{2p} \right) \\ &\leq C^p h \left(h^{2p} \|\nabla V(x)\|^{2p} + d^p h^p p^p \right). \end{aligned}$$

Therefore,

$$\begin{aligned} \mathbb{E}^{\tilde{\mathcal{Q}}_x} \left[\left| \frac{d\tilde{\mathcal{Q}}_x}{d\tilde{\mathcal{Q}}_x} - 1 \right|^p \right] &\leq (C\beta^2 h^2 p)^{p/2} \left(h^{p/2} \|\nabla V(x)\|^p + d^{p/2} p^{p/2} \right) \exp O(\beta^2 h^3 p^2 \|\nabla V(x)\|^2 + \beta^2 d h^2 p^2). \end{aligned}$$

We integrate w.r.t. $\pi(dx)$ and apply Cauchy–Schwarz:

$$\begin{aligned} \int \mathbb{E}^{\tilde{\mathcal{Q}}_x} \left[\left| \frac{d\tilde{\mathcal{Q}}_x}{d\tilde{\mathcal{Q}}_x} - 1 \right|^p \right] \pi(dx) &\leq (C\beta^2 h^2 p)^{p/2} \left(h^{p/2} \sqrt{\int \|\nabla V(x)\|^{2p} \pi(dx)} + d^{p/2} p^{p/2} \right) \\ &\quad \times \int \exp O(\beta^2 h^3 p^2 \|\nabla V(x)\|^2 + \beta^2 d h^2 p^2) \pi(dx). \end{aligned}$$

If we again apply the gradient concentration inequality (Lemma 6.2.7), we obtain

$$\left\{ \int \mathbb{E}^{\tilde{\mathcal{Q}}_x} \left[\left| \frac{d\tilde{\mathcal{Q}}_x}{d\tilde{\mathcal{Q}}_x} - 1 \right|^p \right] \pi(dx) \right\}^{1/p} \lesssim \beta h \sqrt{p} (\sqrt{\beta d h p} + \sqrt{d p}) \lesssim \beta h d^{1/2} p$$

provided that $h \lesssim 1/\beta d^{1/3} p^{2/3}$. By Markov's inequality, it implies that the second term arising from the pointwise projection property is at most $O(\beta d^{1/2} h p / \delta^{1/p})$ with probability at least $1 - \delta$ under π . In particular, if we take $p \asymp \log(1/\delta)$, the bound reads $O(\beta d^{1/2} h \log(1/\delta))$.

Conductance argument.

We now finish with a conductance argument, similarly to Section 7.5. Namely, we must take δ so that $\delta \lesssim s\sqrt{\alpha h} \asymp \varepsilon^2 \sqrt{\alpha h} / \chi^2(\mu_0 \parallel \pi)$ and h so that $h \lesssim 1/(\beta d^{1/2} \log(1/\delta))$. This is satisfied by taking $h \asymp 1/\{\beta \sqrt{d} \log(\chi^2(\mu_0 \parallel \pi)/\varepsilon^2)\}$. Then, Lemma 7.4.8 yields $c_s \gtrsim \sqrt{\alpha h}$. Finally, Theorem 7.4.7 yields mixing in total variation in $\tilde{O}(\kappa \sqrt{d} \log^2(\chi^2(\mu_0 \parallel \pi)/\varepsilon^2))$ iterations.

Lower bound.

Finally, we complement the analysis with a matching lower bound. Under a warm start, Chewi et al. (2021) showed a lower bound of roughly $\tilde{\Omega}(\sqrt{d})$, which was improved in Wu et al. (2021) to $\tilde{\Omega}(\kappa \sqrt{d})$. We state the result here.

Theorem 7.6.5 (Wu et al. (2021)). *For every choice of step size $h > 0$, there exists a target distribution $\pi \propto \exp(-V)$ on $\mathbb{R}^d$ with $I_d \leq \nabla^2 V \leq \kappa I_d$, as well as an initialization μ_0 with $\chi^2(\mu_0 \parallel \pi) \lesssim 1$, such that the number of iterations required for MALA to reach total variation ε from π is at least $\tilde{\Omega}(\kappa \sqrt{d} \log(1/\varepsilon))$.*

Bibliographical Notes

The Metropolis–Hastings algorithm dates back to the birth of MCMC itself (Metropolis et al., 1953; Hastings, 1970), whereas rejection sampling can be traced back to Ulam and von Neumann (von

Neumann, 1951; Eckhardt, 1987). Moreover, MALA and MHMC were introduced in Besag et al. (1995) and Duane et al. (1987) respectively; in the latter reference, the algorithm was introduced as “hybrid Monte Carlo”. The work of Neal (2011) played an important role in popularizing the use of MHMC for Bayesian applications.

Traditionally, these algorithms were studied in an asymptotic framework in which a scaling $h \propto 1/d^\theta$ is sought so that the algorithm admits a non-degenerate diffusion limit as $d \rightarrow \infty$; see Roberts and Rosenthal (1998) for a particularly influential work in this direction. Part of the appeal lies in the precise calculation of the limiting acceptance probabilities, leading to an effective tuning strategy: set the step size of the algorithm so that the proportion of accepted moves equals the limit value.

The quantity d^θ is interpreted as a proxy for the mixing time, and in this way the following predictions were made for the complexities: $\Theta(d^{1/3})$ for MALA, and $\Theta(d^{1/4})$ for MHMC (Beskos et al., 2013). There are notable limitations of the scaling limit approach, however: it imposes stringent assumptions, such as higher-order smoothness or that the target is a product measure, and the analysis is often conducted at stationarity. Indeed, the scalings are different in the non-stationary or “transient” regime (Kuntz et al., 2018; Lee, 2025). There are further works aimed at addressing these issues, but a precise determination of the mixing time requires a non-asymptotic framework. The paper Chewi et al. (2021) further emphasized this point by proving a $\tilde{\Omega}(d^{1/2})$ lower bound for MALA (stated as Theorem 7.6.5), which seemingly contradicts the $\Theta(d^{1/3})$ dependence predicted by the scaling limit. Of course, the contradiction is illusory; the scaling limit requires more assumptions (e.g., higher-order smoothness). Still, this presents a cautionary tale for interpreting asymptotic results with the non-asymptotic, minimax lens that we adopt in this book.

The overlap lemma (Lemma 7.4.6) was refined by a suggestion from Kexin Zhang. The s -conductance mixing time argument given in Exercise 7.5 arose out of discussions with Adil Salim. More generally, it is closely related to the theory of weak Poincaré inequalities; see Andrieu et al. (2022, 2026). The analysis of independent Metropolis–Hastings in Exercise 7.11 is from Liu and Chewi (2026); I thank Sam Power for the observation that $\frac{1}{2}$ -laziness is not needed here.

The $\tilde{O}(\kappa d)$ rate for MALA was first obtained in Dwivedi et al. (2019) with a warm start, and then in Chen et al. (2020) for a cold start. Initially, these rates also had an additional $\tilde{O}(\kappa^{3/2} d^{1/2})$ term, which was removed via a gradient concentration bound in Lee et al. (2020). Our analysis in Section 7.5 is based on Dwivedi et al. (2019), except that we use the gradient bound from Lemma 6.2.7 (which improves, in some respects, the bound used in Chen et al. (2020)). The rate for MRW stated in Theorem 7.3.4 was proven in Andrieu et al. (2024).

The paper Chewi et al. (2021) showed that under standard assumptions, the dimension dependence of MALA improves to $d^{1/2}$ under a warm start. However, due to a typo in the paper, the condition number dependence of the result was unclear. Subsequently, the optimal bound of $\tilde{O}(\kappa d^{1/2})$ was obtained in Wu et al. (2021) via a different approach, which avoids stochastic calculus but relies on rather delicate calculations for the acceptance probability. These papers were later revisited: the rate of Chewi et al. (2021) was clarified to be $\tilde{O}(\kappa d^{1/2} + \kappa^2)$ in the PhD thesis Chewi (2023), and the approach of Wu et al. (2021) was further simplified in Chen and Gatmiry (2023a). In Section 7.6, we follow the approach of Chewi et al. (2021); Chewi (2023), except that we further refine it via Lemma 6.2.7 which to reach the optimal $\tilde{O}(\kappa d^{1/2})$ bound. A warm start for MALA, as in Corollary 7.3.6, was first attained in Altschuler and Chewi (2024a), albeit with a worse dependence on κ .

The following lower bounds for MALA have been obtained: $\tilde{\Omega}(d^{1/2})$ from a warm start (Chewi et al., 2021); $\tilde{\Omega}(\kappa d)$ from a cold start (Lee et al., 2021a); $\tilde{\Omega}(\kappa d^{1/2})$ from a warm start (Wu et al., 2021), improving upon Chewi et al. (2021). Together with the upper bounds, we now have a thorough

understanding of the complexity of MALA for cold vs. warm initializations. We note that [Lee et al. \(2021a\)](#) also contains some lower bounds for MHMC.

There is an extensive literature on variants of these algorithms and analyses which do not fit within the framework considered in this book, and we do not attempt to comprehensively survey it here.

Another approach to designing high-accuracy algorithms based on piecewise deterministic Markov processes (PDMPs). In [Lu and Wang \(2022a\)](#), it is shown that the zigzag sampler can achieve a high-accuracy complexity guarantee of $\tilde{O}(\kappa^2 d^{1/2})$ evaluations of *partial derivatives* of V , from a warm start, in the regime $\kappa \ll d/(\log d)$.

Moreover, in [Chen et al. \(2026b\)](#), it was shown how to simulate rejection sampling using only first-order queries (that is, queries to ∇V). These two approaches show that zeroth-order queries are in fact not necessary for achieving high-accuracy sampling guarantees.

Exercises

An Overview of High-Accuracy Samplers

▷ **Exercise 7.1** (MALA is a special case of MHMC)

Show that when $K = 1$, the MHMC algorithm reduces to MALA.

▷ **Exercise 7.2** (MH filter for the leapfrog integrator)

For the leapfrog integrator F_{leap} , verify that $F_{\text{leap}}^{-1} = R \circ F_{\text{leap}} \circ R$.

Hint: First show that it suffices to consider $T = h$ (i.e., $K = 1$).

Markov Chains in Discrete Time

▷ **Exercise 7.3** (Reversible Markov chains as gradient descent on the Dirichlet energy)

Consider the setting of Section 7.4.

- 1 Show that $\mathcal{E}(f, f) = \frac{1}{2} \mathbb{E}[|f(X_1) - f(X_0)|^2]$, where (X_0, X_1) are two successive iterates of the Markov chain started at stationarity. We also write $\mathcal{E}(f) := \mathcal{E}(f, f)$ as a useful shorthand.
- 2 Show that if $(\mu_n)_{n \in \mathbb{N}}$ are the laws of the iterates of the ℓ -lazy version of P , then the relative densities $(\frac{\mu_n}{\pi})_{n \in \mathbb{N}}$ are the iterates of gradient descent on the Dirichlet energy $\mathcal{E}$ in $L^2(\pi)$. How does the laziness parameter ℓ relate to the step size of the gradient descent?
- 3 Observe that $\mathcal{E}$ is a convex quadratic functional; show that $0 \leq \nabla_{L^2(\pi)}^2 \mathcal{E} \leq 4$. What does the theory of convex optimization suggest for the value of the laziness parameter ℓ ?
- 4 Next, prove a generalization of [Theorem 7.4.2](#) for any value of the laziness parameter $\ell \in [\frac{1}{2}, 1]$ by showing that the spectral gap condition is equivalent to strong convexity of $\mathcal{E}$. Why do we want $\ell \geq \frac{1}{2}$ here?
- 5 Show that the conductance of the chain can also be described as the largest $\mathfrak{c} > 0$ such that for all events $A \subseteq \mathbb{R}^d$, it holds that $\mathcal{E}(\mathbb{1}_A) \geq \mathfrak{c} \|\mathbb{1}_A - \pi(A)\|_{L^2(\pi)}^2$. Hence, conductance can be viewed as a restricted strong convexity condition (restricting the space of functions to indicators of events). In particular, show the bound $\lambda \leq \mathfrak{c}$ in Cheeger's inequality ([Theorem 7.4.3](#)).

▷ **Exercise 7.4** (Entropy decay implies MLSI)

Suppose that P satisfies the following entropy decay condition: there exists $c \in (0, 1)$ such that for all probability measures P ,

$$\text{KL}(\mu P \parallel \pi) \leq (1 - c) \text{KL}(\mu \parallel \pi).$$

Prove that P satisfies a MLSI with constant $C_{\text{MLSI}} \leq 2/c$.

▷ **Exercise 7.5** (Mixing time bound from s -conductance)

In this exercise, we sketch an alternative approach to Theorem 7.4.7.

- 1 Define the following semi-norm: for $s \in (0, 1/2)$,

$$\|f\|_s := \sup_{\pi(A) \leq s} \left| \int f \mathbb{1}_A d\pi \right|.$$

Show that for any $p > 1$, $\|f\|_s \leq s^{1-1/p} \|f\|_{L^p(\pi)}$.

- 2 We define the (s, δ) -spectral gap of P to be the largest $\lambda_{s,\delta} > 0$ such that for all $f \in L^2(\pi)$,

$$\text{var}_\pi f \leq \frac{1}{\lambda_{s,\delta}} \mathcal{E}(f, f) + \delta \|f - \mathbb{E}_\pi f\|_s^2.$$

We typically think of $\delta > 0$ as a universal constant. Show that if $\delta = 1/2$, say, then $\lambda_{s,1/2} \lesssim c_s$. This is the “easy” part of Cheeger’s inequality.

- 3 Using the approach of Exercise 7.3, show that for the iterates $(\mu_n)_{n \in \mathbb{N}}$ of the $\frac{1}{2}$ -lazy chain corresponding to P ,

$$\chi^2(\mu_N \parallel \pi) \leq \left(1 - \frac{\lambda_{s,\delta}}{2}\right)^N \chi^2(\mu_0 \parallel \pi) + \delta \|\rho_0\|_s^2,$$

where $\rho_0 := \frac{d\mu_0}{d\pi} - 1$. Therefore, to prove a mixing time bound under s -conductance, it suffices to establish an analogue of Cheeger’s inequality which lower bounds the (s, δ) -spectral gap in terms of the s -conductance.

- 4 First, we develop a functional version of s -conductance. Show that if s is at most a sufficiently small universal constant, $0 < s \ll 1$, and $\|f - \mathbb{E}_\pi f\|_{2s} \leq 1$, then

$$\mathbb{E}_\pi |f - \mathbb{E}_\pi f| \lesssim \frac{1}{c_s} \mathbb{E} |f(X_1) - f(X_0)| + 1, \quad (7.E.1)$$

where (X_0, X_1) are two successive iterates of the chain started at stationarity.

- 5 Next, show that (7.E.1) implies that for all $f \in L^2(\pi)$,

$$\text{var}_\pi f \lesssim \frac{1}{c_s^2} \mathcal{E}(f, f) + \|(f - \mathbb{E}_\pi f)^2\|_{2s}.$$

Use this to prove a mixing time bound in terms of the s -conductance.

▷ **Exercise 7.6** (s -conductance lemma)

Prove the s -conductance lemma (Lemma 7.4.8).

▷ **Exercise 7.7** (Boosting from TV to χ^2)

The mixing time guarantee in Theorem 7.4.7 is stated for the TV distance, but guarantees for high-accuracy samplers can often be “upgraded” to hold in stronger metrics. To illustrate, suppose that in the setting of Theorem 7.4.7, the Markov chain is initialized at a measure μ_0 with $\frac{d\mu_0}{d\pi} \leq M$. Show that $\chi^2(\mu_N \parallel \pi) \leq \varepsilon^2$ after $N = O((1/c_s^2) \log(M/\varepsilon))$ iterations, where $s = \varepsilon^2/4M^2$.

Hint: Show that $\Delta_s \leq Ms$ and that $\frac{d\mu_n}{d\pi} \leq M$ for all $n \in \mathbb{N}$. Bound $\chi^2(\mu_N \parallel \pi)$ in terms of M and $\|\mu_N - \pi\|_{\text{TV}}$, and then apply Theorem 7.4.7. More generally, this approach shows that if the initialization to the Markov chain satisfies $\mathcal{R}_q(\mu_0 \parallel \pi) < \infty$ for some $q > 1$, then one obtains guarantees for $\mathcal{R}_{q'}(\mu_N \parallel \pi)$ for all $q' < q$.

Analysis of MALA for a Cold Start▷ **Exercise 7.8** (Analysis of MRW)

Follow the analysis in this section and adapt it to the Metropolized random walk (MRW) algorithm. What mixing time bound can you prove?

Analysis of MALA for a Warm Start▷ **Exercise 7.9** (Pointwise projection property)

Prove (7.6.3) from the pointwise projection property.

▷ **Exercise 7.10** (Mixing time for a Gaussian target)

Adapt the analysis in this section to the case when the target distribution is the standard Gaussian. Here, it is possible to do a much more refined analysis; see if you can show that the mixing time of MALA is $\tilde{O}(d^{1/3} \text{polylog}(1/\varepsilon))$ from a warm start. See [Chewi et al. \(2021, Appendix C\)](#) for hints.

▷ **Exercise 7.11** (Independent Metropolis–Hastings)

Consider the independent Metropolis–Hastings algorithm, in which the proposal $Q_x = \mu$ does not depend on x . Check that this defines a positive semidefinite kernel, so that there is no need to make the chain $\frac{1}{2}$ -lazy.

Next, use the analysis in this chapter to prove the following guarantee. Assume that the target π satisfies a Cheeger isoperimetric inequality with some finite constant, and that $\mathcal{R}_q(\mu \parallel \pi) \vee \mathcal{R}_q(\pi \parallel \mu) \leq C_R q^\gamma$, for some $C, C_R, \gamma > 0$ and all $1 \leq q \leq C \log(1/\varepsilon)$. Then, if μ and π are both atomless, $C_R \ll 1/\log^{\gamma+1}(1/\varepsilon)$, and C is sufficiently large, then independent Metropolis–Hastings returns a sample which is ε -close in TV distance to π in $O(\log(1/\varepsilon))$ iterations.

The Proximal Sampler

In this chapter, we discuss the proximal sampler, which was introduced in [Lee et al. \(2021b\)](#) although similar ideas date back earlier. The applications of the proximal sampler include improving the condition number dependence of high-accuracy samplers and extending the sampling guarantees beyond strong log-concavity. Besides these applications, the proximal sampler is interesting in its own right due to its remarkable convergence analysis and its connections with the proximal point method in optimization.

Optimization Box 8.0.1. Given $V : \mathbb{R}^d \rightarrow \mathbb{R}$, the **proximal oracle** associated with V and step size $h > 0$ is the mapping $\text{prox}_{hV} : \mathbb{R}^d \rightarrow \mathbb{R}^d$ given by $\text{prox}_{hV}(x) := \arg \min_{y \in \mathbb{R}^d} \{V(y) + \frac{1}{2h} \|y - x\|^2\}$.

Implementing the proximal oracle. Note that the proximal oracle can be viewed as a regularized version of the original optimization problem. In particular, if $\|\nabla^2 V\|_{\text{op}} \leq \beta$ and $h \leq 1/2\beta$, then the proximal oracle is the solution to a $(\frac{1}{h} - \beta)$ -strongly convex and $(\frac{1}{h} + \beta)$ -smooth optimization problem, with condition number $\frac{1/h + \beta}{1/h - \beta} \leq 3$. Therefore, it can be approximately implemented via convex optimization.

Using the proximal oracle. Given access to the proximal oracle for V , the simplest way to use it to minimize V is to iterate it: $x_{n+1} := \text{prox}_{hV}(x_n)$. This is known as the **proximal point method (PPM)**. If V is α -strongly convex, then we can obtain a convergence rate for the PPM as follows. First, note that the first-order optimality condition for the PPM can be written $\text{prox}_{hV}(x) = (\text{id} + h \nabla V)^{-1}(x)$. From this, one can show that prox_{hV} is $1/(1 + \alpha h)$ -Lipschitz, and $\text{prox}_{hV}(x_\star) = x_\star$. Therefore, $\|x_N - x_\star\| \leq (1 + \alpha h)^{-N} \|x_0 - x_\star\|$. Note that unlike gradient descent, the convergence of the PPM does *not* require smoothness of V .

Application to composite optimization. Arguably, the most well-known application of the proximal oracle is for composite minimization, see [Optimization Box 8.6.4](#).

The PPM can be written $x_{n+1} = x_n - h \nabla V(x_{n+1})$ which shows that it is an *implicit* discretization of the gradient flow. Implicit discretizations are more stable (indeed, note that the convergence rate for the PPM given above holds for *all* choices of $h > 0$, whereas GD typically requires $h < 2/\beta$), but at the cost of being harder to implement. We will see that similar intuitions hold in the context of sampling.

8.1 Introduction to the Proximal Sampler

Let $\pi \propto \exp(-V)$ denote the target distribution. We fix $h > 0$ and define the augmented target distribution

$$\pi(x, y) \propto \exp\left(-V(x) - \frac{\|y - x\|^2}{2h}\right).$$

To avoid confusion, we will explicitly write $\pi^X = \pi$ for the X -marginal, and π^Y for the Y -marginal. Similarly, $\pi^{X|Y}$ and $\pi^{Y|X}$ denote the conditional distributions.

The proximal sampler applies Gibbs sampling to the augmented target. Explicitly, the updates of the proximal sampler are as follows.

Proximal Sampler: Initialize $X_0 \sim \mu_0$. For $n = 0, 1, 2, \dots$:

- 1 Draw $Y_n \sim \pi^{Y|X}(\cdot | X_n) = \text{normal}(X_n, hI_d)$.
- 2 Draw $X_{n+1} \sim \pi^{X|Y}(\cdot | Y_n)$.

Since Gibbs sampling always forms a reversible Markov chain with respect to the target distribution, we conclude that the proximal sampler is *unbiased*: its stationary distribution of the proximal sampler is π . As written, however, the proximal sampler is an idealized algorithm because it is not yet clear how to implement the second step of sampling from $\pi^{X|Y}$. Note that

$$\pi^{X|Y}(x | y) \propto_x \exp\left(-V(x) - \frac{\|y - x\|^2}{2h}\right).$$

Also, recall that in optimization, we wish to minimize the function V , whereas in sampling we want to sample from $\pi \propto \exp(-V)$. Via this correspondence, we see that the optimization analogue of sampling from $\pi^{X|Y}$, which is known as the **restricted Gaussian oracle (RGO)**, is precisely the computation of the proximal oracle. See [Exercise 8.1](#) for another connection between the proximal sampler and the PPM.

Interpretation of the proximal sampler.

At this stage, it may be unclear why the proximal sampler should be interpreted as the PPM for sampling. Part of the justification lies in the convergence results we establish in subsequent sections that are exactly analogous to optimization guarantees for the PPM.

Here, we present another justification in the spirit of the discussion from [Section 4.3](#). There, we observed that sampling, which can be viewed as the minimization of the KL divergence $\text{KL}(\cdot \| \pi)$, is a composite optimization problem over the space of probability measures consisting of a smooth term (the “energy”) and a non-smooth term (the “entropy”). The optimization wisdom from [Optimization Box 8.0.1](#) suggests considering a proximal gradient method in the Wasserstein space, which would require computation of $\text{prox}_{h\mathcal{H}}(\mu) := \arg \min_{\nu \in \mathcal{P}_2(\mathbb{R}^d)} \{\mathcal{H}(\nu) + \frac{1}{2} W_2^2(\mu, \nu)\}$. Such a scheme can be shown to enjoy strong convergence guarantees—see [Salim et al. \(2020\)](#)—but unfortunately the proximal oracle for the entropy functional is not easily implementable.

Recall from [Section 4.3](#) that LMC is a splitting scheme which applies GD to the energy, and *gradient flow* to the entropy, but that this scheme is biased. The paper [Salim et al. \(2020\)](#) shows that if we keep the GD step for the energy, then the “natural” (but intractable) method to apply to the entropy to make the whole algorithm unbiased is the proximal oracle $\text{prox}_{h\mathcal{H}}$. On the other hand, what if we keep the gradient flow for the entropy—i.e., the heat flow—and ask what method we can apply to the energy to keep the algorithm unbiased? In a sense, the proximal sampler is the answer to this

question, since the first step of sampling from a Gaussian is indeed the heat flow. See [Exercise 8.2](#) for an interpretation of the RGO as the *adjoint* to the heat flow in a precise sense.

This discussion highlights that even when V is smooth, there are potential advantages to be gained by consideration of proximal methodology, since the sampling problem carries with it an *inherent* non-smoothness from the entropy functional.

Implementability of the RGO.

In order to obtain an actual algorithm from the proximal sampler, an implementation of the RGO must be provided. Obviously, we have

$$\boxed{\text{complexity of the proximal sampler}} = \boxed{\text{number of iterations of the proximal sampler}} \times \boxed{\text{complexity of implementing the RGO}}.$$

The next few sections will be devoted to studying the iteration complexity for the proximal sampler under various assumptions, including strong log-concavity and functional inequalities such as log-Sobolev inequality.

Implementations of the RGO will be discussed in Section 8.6. Similarly as in [Optimization Box 8.0.1](#), as soon as $\|\nabla^2 V\|_{\text{op}} \leq \beta$ and $h \leq 1/2\beta$, the RGO is a strongly log-concave and log-smooth distribution with condition number at most 3. Therefore, the cost of implementing the RGO will boil down to the cost of sampling from *well-conditioned strongly log-concave* distributions.

Notation.

We write μ_n^X for the law of X_n and μ_n^Y for the law of Y_n for the iterates of the proximal sampler. Observe that if $(Q_t)_{t \geq 0}$ denotes the standard heat semigroup, i.e. $\mu Q_t = \mu * \text{normal}(0, tI_d)$, then $\mu_n^Y = \mu_n^X Q_h$ and $\pi^Y = \pi^X Q_h$.

We also abbreviate $\pi^{X|Y}(\cdot | y)$ as $\pi^{X|Y=y}$.

8.2 Convergence under Strong Log-Concavity

One of the most remarkable features of the proximal sampler is that its convergence analysis closely mirrors the continuous-time theory for the Langevin diffusion. In this section, we initiate this study starting with the strongly log-concave case.

Recall that under strong log-concavity, we have contraction of the Langevin diffusion ([Theorem 1.4.12](#)). We prove the analogue of this fact for the proximal sampler.

Theorem 8.2.1 (Contractivity of the proximal sampler). *Assume that the target π^X is α -strongly log-concave. Also, let $(\mu_n^X)_{n \in \mathbb{N}}$ and $(\bar{\mu}_n^X)_{n \in \mathbb{N}}$ denote two runs of the proximal sampler with target π^X . Then,*

$$W_2(\mu_n^X, \bar{\mu}_n^X) \leq \frac{W_2(\mu_0^X, \bar{\mu}_0^X)}{(1 + \alpha h)^n}.$$

The contraction factor matches the contraction for the proximal point method in optimization, see [Exercise 8.3](#). Since π^X is left invariant by the proximal sampler, the contraction result also implies a convergence result in W_2 .

We will give two proofs of this theorem. First, note that it suffices to consider one iteration and to

prove $W_2(\mu_1^X, \bar{\mu}_1^X) \leq \frac{1}{1+\alpha h} W_2(\mu_0^X, \bar{\mu}_0^X)$. Next, since the heat flow is a Wasserstein contraction (which follows from [Theorem 1.4.12](#) but can also be proven by a straightforward coupling), it holds that $W_2(\mu_0^Y, \bar{\mu}_0^Y) \leq W_2(\mu_0^X, \bar{\mu}_0^X)$, so it suffices to show $W_2(\mu_1^X, \bar{\mu}_1^X) \leq \frac{1}{1+\alpha h} W_2(\mu_0^Y, \bar{\mu}_0^Y)$.

We will use the following coupling lemma.

Lemma 8.2.2. *Suppose that for all $y, \bar{y} \in \mathbb{R}^d$, we have*

$$W_2(\pi^{X|Y=y}, \pi^{X|Y=\bar{y}}) \leq C \|y - \bar{y}\|. \quad (8.2.3)$$

Then, $W_2(\mu_1^X, \bar{\mu}_1^X) \leq C W_2(\mu_0^Y, \bar{\mu}_0^Y)$.

The intuition is that since μ_1^X and $\bar{\mu}_1^X$ are obtained from μ_0^Y and $\bar{\mu}_0^Y$ by sampling from the RGO $\pi^{X|Y}$, the contraction statement in (8.2.3) can be used to bound $W_2(\mu_1^X, \bar{\mu}_1^X)$. The proof of the lemma is relatively straightforward and good practice for working with couplings, so it is left as [Exercise 8.4](#).

The first proof we present is from [Lee et al. \(2021c\)](#).

Proof of Theorem 8.2.1 via functional inequalities To prove (8.2.3), we note that $\pi^{X|Y}(\cdot | \bar{y})$ is $(\alpha + \frac{1}{h})$ -strongly log-concave. Recall that by the Bakry–Émery theorem ([Theorem 1.2.30](#)) and the Otto–Villani theorem ([Exercise 1.16](#)) this implies the log-Sobolev inequality (1.4.15) and Talagrand’s T_2 inequality (1.4.16). Applying these inequalities,

$$W_2^2(\pi^{X|Y=y}, \pi^{X|Y=\bar{y}}) \leq \frac{2}{\alpha + \frac{1}{h}} \text{KL}(\pi^{X|Y=y} \parallel \pi^{X|Y=\bar{y}}) \leq \frac{1}{(\alpha + \frac{1}{h})^2} \text{FI}(\pi^{X|Y=y} \parallel \pi^{X|Y=\bar{y}}).$$

We can compute the Fisher information explicitly. Indeed,

$$\nabla \log \frac{\pi^{X|Y=y}}{\pi^{X|Y=\bar{y}}} = \nabla \left(\frac{\|\bar{y} - \cdot\|^2}{2h} - \frac{\|y - \cdot\|^2}{2h} \right) = \frac{y - \bar{y}}{h}$$

so that

$$\text{FI}(\pi^{X|Y=y} \parallel \pi^{X|Y=\bar{y}}) = \mathbb{E}_{\pi^{X|Y=y}} \left[\left\| \nabla \log \frac{\pi^{X|Y=y}}{\pi^{X|Y=\bar{y}}} \right\|^2 \right] = \frac{\|y - \bar{y}\|^2}{h^2}.$$

Hence,

$$W_2^2(\pi^{X|Y=y}, \pi^{X|Y=\bar{y}}) \leq \frac{1}{(\alpha + \frac{1}{h})^2} \frac{\|y - \bar{y}\|^2}{h^2} = \frac{1}{(1 + \alpha h)^2} \|y - \bar{y}\|^2. \quad \square$$

The next proof, from [Chen et al. \(2022b\)](#), directly uses strong convexity in Wasserstein space.

Proof of Theorem 8.2.1 via Wasserstein calculus This proof rests on the following interpretation of the RGO. Let $\mathcal{F}(\mu) := \text{KL}(\mu \parallel \pi^X)$. Then, by [Exercise 8.1](#),

$$\pi^{X|Y=y} = \arg \min_{\mu \in \mathcal{P}_2(\mathbb{R}^d)} \left\{ \mathcal{F}(\mu) + \frac{1}{2h} W_2^2(\mu, \delta_y) \right\} =: \text{prox}_{h\mathcal{F}}(\delta_y).$$

The first-order optimality condition on the Wasserstein space ([Ambrosio et al., 2008](#), Lemma 10.1.2) reads

$$0 \in \partial \mathcal{F}(\pi^{X|Y=y}) + \frac{1}{h} (\text{id} - y), \quad \pi^{X|Y=y}\text{-a.s.}$$

where $\partial \mathcal{F}$ is the subdifferential of $\mathcal{F}$ on Wasserstein space.

Using this, we obtain

$$\begin{aligned} \text{id} &\in y - h \partial \mathcal{F}(\pi^{X|Y=y}), & \pi^{X|Y=y}\text{-a.s.} \\ \text{id} &\in \bar{y} - h \partial \mathcal{F}(\pi^{X|Y=\bar{y}}), & \pi^{X|Y=\bar{y}}\text{-a.s.} \end{aligned}$$

Let T be the optimal transport map from $\pi^{X|Y=y}$ to $\pi^{X|Y=\bar{y}}$. The second condition above can then be rewritten as

$$T \in \bar{y} - h \partial \mathcal{F}(\pi^{X|Y=\bar{y}}) \circ T, \quad \pi^{X|Y=y}\text{-a.s.}$$

We now abuse notation and write $\partial \mathcal{F}(\pi^{X|Y=y})$ for a particular element of the subdifferential and similarly for $\partial \mathcal{F}(\pi^{X|Y=\bar{y}})$. Then, $\pi^{X|Y=y}\text{-a.s.}$,

$$\begin{aligned} \|T - \text{id}\|^2 &= \|\bar{y} - y\|^2 - 2h \langle \partial \mathcal{F}(\pi^{X|Y=\bar{y}}) \circ T - \partial \mathcal{F}(\pi^{X|Y=y}), T - \text{id} \rangle \\ &\quad - h^2 \|\partial \mathcal{F}(\pi^{X|Y=\bar{y}}) \circ T - \partial \mathcal{F}(\pi^{X|Y=y})\|^2. \end{aligned}$$

We now integrate w.r.t. $\pi^{X|Y=y}$ and apply strong convexity of $\mathcal{F}$ in Wasserstein space:

$$W_2^2(\pi^{X|Y=y}, \pi^{X|Y=\bar{y}}) \leq \|y - \bar{y}\|^2 - 2\alpha h W_2^2(\pi^{X|Y=y}, \pi^{X|Y=\bar{y}}) - \alpha^2 h^2 W_2^2(\pi^{X|Y=y}, \pi^{X|Y=\bar{y}})$$

and hence

$$W_2^2(\pi^{X|Y=y}, \pi^{X|Y=\bar{y}}) \leq \frac{1}{(1 + \alpha h)^2} \|y - \bar{y}\|^2. \quad \square$$

The point of the second proof is that, although it uses some heavy machinery, it is just a translation of a Euclidean optimization proof into the language of Wasserstein space (see [Exercise 8.3](#)).

8.3 Simultaneous Flow and Time Reversal

We now introduce two new techniques in order to further analyze the proximal sampler.

Simultaneous flow.

The first technique is based on the observation that in going from μ_n^X to μ_n^Y , and from π^X to π^Y , we are applying the heat flow. Given any f -divergence $\mathcal{D}_f(\cdot \| \cdot)$, we will compute its time derivative when both arguments undergo simultaneous heat flow. Remarkably, the result will be almost the same as the time derivative of the f -divergence to the target along the continuous-time Langevin diffusion, in a sense to be made precise. The upshot is that the analysis of the proximal sampler closely resembles the analysis of the continuous-time Langevin diffusion.

The simultaneous heat flow calculation is inspired by [Vempala and Wibisono \(2019\)](#), and was carried out for the proximal sampler in [Chen et al. \(2022b\)](#). Here, we place the calculation in the framework of abstract Markov semigroup theory in order to show that the final result is a consequence of general principles rather than surprising coincidences.

Let $f : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ be a convex function with $f(1) = 0$, and let $\mathcal{D}_f$ be the associated f -divergence (see [Section 1.5](#)). We prove the following abstract result.

Theorem 8.3.1 (Simultaneous flow). *Suppose that we have two processes evolving via*

$$\partial_t \mu_t = \mathcal{L}_t^* \mu_t, \quad \partial_t \nu_t = \mathcal{L}_t^* \nu_t,$$

where $$ denotes the adjoint with respect to some reference measure $\mathfrak{m}$. Let Γ_t denote the carré du champ operator associated with $\mathcal{L}_t$,*

$$\Gamma_t(f, g) := \frac{1}{2} (\mathcal{L}_t(fg) - f \mathcal{L}_t g - g \mathcal{L}_t f), \quad \Gamma_t(f) := \Gamma_t(f, f).$$

Assume that $\mathcal{L}_t$ satisfies the diffusion chain rule (Definition 2.2.13) for all $t \geq 0$. Then,

$$\partial_t \mathcal{D}_f(\mu_t \parallel \nu_t) = - \int f'' \left(\frac{d\mu_t}{d\nu_t} \right) \Gamma_t \left(\frac{d\mu_t}{d\nu_t} \right) d\nu_t.$$

Proof We identify μ_t and ν_t with their densities w.r.t. $\mathfrak{m}$, and all of the following integrals are w.r.t. $\mathfrak{m}$ (although we do not write this explicitly to keep the notation light). By the definition of the adjoint,

$$\begin{aligned} \partial_t \int f \left(\frac{\mu_t}{\nu_t} \right) \nu_t &= \int f' \left(\frac{\mu_t}{\nu_t} \right) (\partial_t \mu_t - \frac{\mu_t}{\nu_t} \partial_t \nu_t) + \int f \left(\frac{\mu_t}{\nu_t} \right) \partial_t \nu_t \\ &= \int f' \left(\frac{\mu_t}{\nu_t} \right) (\mathcal{L}_t^* \mu_t - \frac{\mu_t}{\nu_t} \mathcal{L}_t^* \nu_t) + \int f \left(\frac{\mu_t}{\nu_t} \right) \mathcal{L}_t^* \nu_t \\ &= \int \left(\mathcal{L}_t f' \left(\frac{\mu_t}{\nu_t} \right) \mu_t - \mathcal{L}_t \left(f' \left(\frac{\mu_t}{\nu_t} \right) \frac{\mu_t}{\nu_t} \right) \nu_t + \mathcal{L}_t f \left(\frac{\mu_t}{\nu_t} \right) \nu_t \right). \end{aligned}$$

For the second term, we apply the definition of the carré du champ to obtain

$$\mathcal{L}_t \left(f' \left(\frac{\mu_t}{\nu_t} \right) \frac{\mu_t}{\nu_t} \right) \nu_t = 2 \Gamma_t \left(f' \left(\frac{\mu_t}{\nu_t} \right), \frac{\mu_t}{\nu_t} \right) \nu_t + \mu_t \mathcal{L}_t f' \left(\frac{\mu_t}{\nu_t} \right) + f' \left(\frac{\mu_t}{\nu_t} \right) \nu_t \mathcal{L}_t \frac{\mu_t}{\nu_t}.$$

For the third term, the diffusion chain rule says that for any function f ,

$$\mathcal{L}_t f(\rho) = f'(\rho) \mathcal{L}_t \rho + f''(\rho) \Gamma_t(\rho)$$

hence

$$\mathcal{L}_t f \left(\frac{\mu_t}{\nu_t} \right) = f' \left(\frac{\mu_t}{\nu_t} \right) \mathcal{L}_t \frac{\mu_t}{\nu_t} + f'' \left(\frac{\mu_t}{\nu_t} \right) \Gamma_t \left(\frac{\mu_t}{\nu_t} \right).$$

Substituting this in and taking advantage of a cancellation,

$$\partial_t \int f \left(\frac{\mu_t}{\nu_t} \right) \nu_t = -2 \int \Gamma_t \left(f' \left(\frac{\mu_t}{\nu_t} \right), \frac{\mu_t}{\nu_t} \right) \nu_t + \int f'' \left(\frac{\mu_t}{\nu_t} \right) \Gamma_t \left(\frac{\mu_t}{\nu_t} \right) \nu_t.$$

Finally, to simplify this further, the diffusion chain rule implies

$$\Gamma_t(f(\rho), \sigma) = f'(\rho) \Gamma_t(\rho, \sigma),$$

and it finishes the proof. $\square$

Let us now parse the content of the theorem in several important cases. First, suppose that $(\mu_t)_{t \geq 0}$ is the Langevin diffusion and $\nu_t = \pi$ is the stationary distribution for all $t \geq 0$. In this case, $\mathcal{L}_t = \mathcal{L}$ is the Langevin generator for all $t \geq 0$, $\Gamma_t(\rho) = \|\nabla \rho\|^2$, and we obtain the following formula: $\partial_t \mathcal{D}_f(\mu_t \parallel \pi) = - \int f''(\mu_t/\pi) \|\nabla(\mu_t/\pi)\|^2 d\pi$. In particular, for $f(x) = x \log x$, it recovers $\partial_t \text{KL}(\mu_t \parallel \pi) = -\text{FI}(\mu_t \parallel \pi)$.

Now suppose that $(\mu_t)_{t \geq 0}, (\nu_t)_{t \geq 0}$ are two copies of the Langevin diffusion, where $(\nu_t)_{t \geq 0}$ is not necessarily stationary. Then, we can obtain the formula $\partial_t \mathcal{D}_f(\mu_t \parallel \nu_t) = - \int f''(\mu_t/\nu_t) \|\nabla(\mu_t/\nu_t)\|^2 d\nu_t$,

which contains $\partial_t \text{KL}(\mu_t \parallel \nu_t) = -\text{FI}(\mu_t \parallel \nu_t)$ as a special case. We see that the evolution of the KL divergence as *both* arguments evolve is hardly any more complicated than the case where only *one* of the arguments evolves!

A further special case is obtained when $(\mu_t)_{t \geq 0}, (\nu_t)_{t \geq 0}$ evolve by the heat flow, which corresponds to the Langevin diffusion with potential $V = 0$ (up to time rescaling). The theorem still applies, we obtain $\partial_t \text{KL}(\mu_t \parallel \nu_t) = -\frac{1}{2} \text{FI}(\mu_t \parallel \nu_t)$ (the factor of $\frac{1}{2}$ arises because the generator for the heat flow is $\mathcal{L} = \frac{1}{2} \Delta$, which then yields $\Gamma = \frac{1}{2} \|\nabla \cdot\|^2$). Note that [Theorem 8.3.1](#) explains why: the formula for the derivative in time of the f -divergence only depends on the diffusion through its *carré du champ*, which is the same (up to rescaling) for both the heat flow and for the Langevin diffusion. This shows that the heat flow is not really that special and that much of the following discussion on the proximal sampler would equally well apply if we replaced the heat flow with, e.g., the OU process, but the heat flow is convenient for both calculation and implementation.

Although this theorem is already enough to prove new convergence results for the proximal sampler, the rates will be slightly suboptimal. The reason for this is because we have only considered one step of the proximal sampler, in which the algorithm goes from μ_n^X to μ_n^Y (and the target goes from π^X to π^Y). In order to obtain the sharp convergence rates, we also need to consider the second step, in which we go from μ_n^Y to μ_{n+1}^X (and the target returns from π^Y to π^X). For reasons that will become clear shortly, we refer to these steps as the “forward step” and the “backward step” respectively.

First, consider the evolution of the target along the heat semigroup $t \mapsto \pi^X Q_t$, so that at time h we arrive at π^Y . The stochastic process representation of this evolution is $dZ_t = dB_t$, with $Z_0 \sim \pi^X$ and $Z_h \sim \pi^Y$, thus describing the forward step. So far so good, but how should we think about the backward step? By definition, π^X is obtained from π^Y by the relation $\pi^X = \int \pi^{X|Y=y} \pi^Y(dy)$, but this is not as helpful because we lose the stochastic process view which allows us to apply calculus. Instead, we will think of π^X as being obtained from π^Y by the *time reversal* of the diffusion $(Z_t)_{t \in [0, h]}$.

Time reversal and simultaneous backward flow.

We now apply the result of [Section 3.3.2](#), which implies that the time reversal of the SDE

$$dZ_t = dB_t, \quad Z_0 \sim \pi^X$$

is given by the SDE

$$dZ_t^\leftarrow = \nabla \log(\pi^X Q_{h-t})(Z_t^\leftarrow) dt + dB_t \quad (8.3.2)$$

in the sense that if we initialize the SDE at $Z_0^\leftarrow = y$, then $Z_h^\leftarrow \sim \pi^{X|Y=y}$.

In particular, if we initialize the process at $Z_0^\leftarrow \sim \pi^Y$, then $Z_h^\leftarrow \sim \pi^X$. On the other hand, if we initialize the process at $Z_0^\leftarrow \sim \mu_n^Y$, then the law of $Z_h^\leftarrow$ is $\int \pi^{X|Y=y} \mu_n^Y(dy) = \mu_{n+1}^X$. Thus, we have successfully exhibited a stochastic process representation which takes us from μ_n^Y to μ_{n+1}^X .

For any measure μ , write $\mu_t^\leftarrow := \mu Q_t^\leftarrow$ for the law of $Z_t^\leftarrow$ initialized at $Z_0^\leftarrow \sim \mu$. If we have two such processes $(\mu_t^\leftarrow)_{t \in [0, h]}, (\nu_t^\leftarrow)_{t \in [0, h]}$, then instantaneously at time t they evolve according a generator $\mathcal{L}_t$, which corresponds to the carré du champ $\Gamma_t(\rho) = \frac{1}{2} \|\nabla \rho\|^2$. So once again, by [Theorem 8.3.1](#) we obtain $\partial_t \text{KL}(\mu_t^\leftarrow \parallel \nu_t^\leftarrow) = -\frac{1}{2} \text{FI}(\mu_t^\leftarrow \parallel \nu_t^\leftarrow)$, which leads to a pleasing symmetry between the forward and backward steps of the proximal sampler.

8.4 Convergence under Log-Concavity

Next, we present a convergence proof for the proximal sampler under log-concavity, following [Chen et al. \(2022b\)](#). The proof can be compared to the $1/t$ convergence rate for the Langevin diffusion under log-concavity (1.4.14), which was obtained via a Lyapunov function argument.

Theorem 8.4.1 (Proximal sampler under log-concavity). *Assume that the target π^X is log-concave. Then, for the law μ_n^X of the n -th iterate of the proximal sampler,*

$$\text{KL}(\mu_n^X \parallel \pi^X) \leq \frac{W_2^2(\mu_0^X, \pi^X)}{nh}.$$

Proof **Forward step.** Along the simultaneous heat flow, [Theorem 8.3.1](#) shows that

$$\partial_t \text{KL}(\mu_0^X Q_t \parallel \pi^X Q_t) = -\frac{1}{2} \text{FI}(\mu_0^X Q_t \parallel \pi^X Q_t)$$

so we need to lower bound the Fisher information. Also, log-concavity is preserved by convolution, so $\pi^X Q_t$ is log-concave ([Exercise 2.11](#)). Hence, by convexity of the KL divergence to a log-concave target along Wasserstein geodesics ([Theorem 1.4.5](#)),

$$0 = \text{KL}(\pi^X Q_t \parallel \pi^X Q_t) \geq \text{KL}(\mu_0^X Q_t \parallel \pi^X Q_t) + \mathbb{E}_{\mu_0^X Q_t} \left\langle \nabla \log \frac{\mu_0^X Q_t}{\pi^X Q_t}, T_{\mu_0^X Q_t \rightarrow \pi^X Q_t} - \text{id} \right\rangle.$$

Rearranging this and using the Cauchy–Schwarz inequality,

$$\underbrace{\mathbb{E}_{\mu_0^X Q_t} \left[\left\| \nabla \log \frac{\mu_0^X Q_t}{\pi^X Q_t} \right\|^2 \right]}_{\text{FI}(\mu_0^X Q_t \parallel \pi^X Q_t)} W_2^2(\mu_0^X Q_t, \pi^X Q_t) \geq \text{KL}(\mu_0^X Q_t \parallel \pi^X Q_t)^2.$$

Combining this with the fact that the Wasserstein distance decreases along the simultaneous heat flow,

$$\partial_t \text{KL}(\mu_0^X Q_t \parallel \pi^X Q_t) \leq -\frac{1}{2} \frac{\text{KL}(\mu_0^X Q_t \parallel \pi^X Q_t)^2}{W_2^2(\mu_0^X, \pi^X)}.$$

Solving this differential inequality,

$$\frac{1}{\text{KL}(\mu_0^Y \parallel \pi^Y)} = \frac{1}{\text{KL}(\mu_0^X Q_h \parallel \pi^X Q_h)} \geq \frac{1}{\text{KL}(\mu_0^X \parallel \pi^X)} + \frac{h}{2 W_2^2(\mu_0^X, \pi^X)}.$$

Backward step. Along the simultaneous backward heat flow, [Theorem 8.3.1](#) gives

$$\partial_t \text{KL}(\mu_0^Y Q_t^\leftarrow \parallel \pi^Y Q_t^\leftarrow) = -\frac{1}{2} \text{FI}(\mu_0^Y Q_t^\leftarrow \parallel \pi^Y Q_t^\leftarrow).$$

Since $\pi^Y Q_t^\leftarrow = \pi^X Q_{h-t}$ is log-concave and $t \mapsto W_2^2(\mu_0^Y Q_t^\leftarrow, \pi^Y Q_t^\leftarrow)$ is decreasing (which is checked via a coupling argument using the diffusion (8.3.2)), a similar calculation as the forward step leads to the inequality

$$\frac{1}{\text{KL}(\mu_1^X \parallel \pi^X)} = \frac{1}{\text{KL}(\mu_0^Y Q_h^\leftarrow \parallel \pi^Y Q_h^\leftarrow)} \geq \frac{1}{\text{KL}(\mu_0^Y \parallel \pi^Y)} + \frac{h}{2 W_2^2(\mu_0^X, \pi^X)}.$$

We iterate these inequalities, using the fact that $W_2(\mu_n^X, \pi^X) \leq W_2(\mu_0^X, \pi^X)$ for all $n \in \mathbb{N}$ (which follows from [Theorem 8.2.1](#)) to obtain

$$\frac{1}{\text{KL}(\mu_n^X \parallel \pi^X)} \geq \frac{1}{\text{KL}(\mu_0^X \parallel \pi^X)} + \frac{nh}{W_2^2(\mu_0^X, \pi^X)}$$

or

$$\text{KL}(\mu_n^X \parallel \pi^X) \leq \frac{\text{KL}(\mu_0^X \parallel \pi^X)}{1 + nh \text{KL}(\mu_0^X \parallel \pi^X)/W_2^2(\mu_0^X, \pi^X)} \leq \frac{W_2^2(\mu_0^X, \pi^X)}{nh}. \quad \square$$

8.5 Convergence under Functional Inequalities

We now prove convergence guarantees for the proximal sampler when the target satisfies either a Poincaré inequality or a log-Sobolev inequality, following [Chen et al. \(2022b\)](#).

Theorem 8.5.1 (Proximal sampler under PI). *Suppose that the target π^X satisfies a Poincaré inequality with constant C_{PI} . Then, for the law μ_n^X of the n -th iterate of the proximal sampler,*

$$\chi^2(\mu_n^X \parallel \pi^X) \leq \frac{\chi^2(\mu_0^X \parallel \pi^X)}{(1 + h/C_{\text{PI}})^{2n}}.$$

Proof **Forward step.** Along the simultaneous heat flow, by [Theorem 8.3.1](#),

$$\partial_t \chi^2(\mu_0^X Q_t \parallel \pi^X Q_t) = -\mathbb{E}_{\pi^X Q_t} \left[\left\| \nabla \frac{\mu_0^X Q_t}{\pi^X Q_t} \right\|^2 \right].$$

Since π^X satisfies a Poincaré inequality with constant C_{PI} , by subadditivity of the Poincaré constant under convolution ([Proposition 2.3.8](#)), $\pi^X Q_t$ satisfies a Poincaré inequality with constant at most $C_{\text{PI}} + t$. It therefore yields

$$\partial_t \chi^2(\mu_0^X Q_t \parallel \pi^X Q_t) \leq -\frac{1}{C_{\text{PI}} + t} \text{var}_{\pi^X Q_t} \frac{\mu_0^X Q_t}{\pi^X Q_t} = -\frac{1}{C_{\text{PI}} + t} \chi^2(\mu_0^X Q_t \parallel \pi^X Q_t)$$

and hence

$$\chi^2(\mu_0^Y \parallel \pi^Y) = \chi^2(\mu_0^X Q_h \parallel \pi^X Q_h) \leq \exp\left(-\int_0^h \frac{1}{C_{\text{PI}} + t} dt\right) \chi^2(\mu_0^X \parallel \pi^X) = \frac{\chi^2(\mu_0^X \parallel \pi^X)}{1 + h/C_{\text{PI}}}.$$

Backward step. Along the simultaneous backward heat flow, [Theorem 8.3.1](#) yields

$$\partial_t \chi^2(\mu_0^Y Q_t^{\leftarrow} \parallel \pi^Y Q_t^{\leftarrow}) = -\mathbb{E}_{\pi^Y Q_t^{\leftarrow}} \left[\left\| \nabla \frac{\mu_0^Y Q_t^{\leftarrow}}{\pi^Y Q_t^{\leftarrow}} \right\|^2 \right].$$

Using the fact that $\pi^Y Q_t^{\leftarrow} = \pi^X Q_{h-t}$ satisfies the Poincaré inequality with constant at most $C_{\text{PI}} + h - t$, we deduce similarly that

$$\chi^2(\mu_1^X \parallel \pi^X) = \chi^2(\mu_0^Y Q_h^{\leftarrow} \parallel \pi^Y Q_h^{\leftarrow}) \leq \frac{\chi^2(\mu_0^Y \parallel \pi^Y)}{1 + h/C_{\text{PI}}}.$$

Iterating this pair of inequalities yields the result. $\square$

A similar result holds for the log-Sobolev inequality; since the proof is entirely analogous, we leave it as [Exercise 8.5](#).

Theorem 8.5.2 (Proximal sampler under LSI). *Suppose that the target π^X satisfies a log-Sobolev inequality with constant C_{LSI} . Then, for the law μ_n^X of the n -th iterate of the proximal sampler,*

$$\text{KL}(\mu_n^X \parallel \pi^X) \leq \frac{\text{KL}(\mu_0^X \parallel \pi^X)}{(1 + h/C_{\text{LSI}})^{2n}}.$$

Recall that if π^X is α -strongly log-concave, then it satisfies a log-Sobolev inequality with constant $C_{\text{LSI}} \leq 1/\alpha$ (Theorem 1.2.30). Thus, the contraction factor of $\frac{1}{(1+\alpha h)^2}$ in KL divergence matches the contraction factor in W_2^2 distance (Theorem 8.2.1). To get this sharp result, it is necessary to utilize the backward step.

Similarly to Theorem 2.2.25, it is also possible to obtain guarantees for Rényi divergences, see Exercise 8.6.

Remark 8.5.3. It is a curious observation that in the W_2 guarantee of Theorem 8.2.1, the contraction factor of $\frac{1}{(1+\alpha h)^2}$ occurs solely in the backward step, whereas in Theorem 8.5.2 the forward and backward steps each contribute a contraction factor of $\frac{1}{1+\alpha h}$.

8.6 Implementations of the RGO and Applications

In this section, we discuss various implementation strategies for the RGO, thereby obtaining concrete sampling guarantees.

Implementation via naïve rejection sampling.

First, we consider a simple implementation of the RGO based on rejection sampling, which we studied in Section 7.1. Suppose that the potential V is β -smooth. Recall that for $V_y(x) := V(x) + \frac{1}{2h} \|y - x\|^2$, we have $(\frac{1}{h} - \beta) I_d \leq \nabla^2 V_y \leq (\frac{1}{h} + \beta) I_d$. In particular, for $h < 1/\beta$, the condition number of V_y is $\kappa = (\frac{1}{h} + \beta)/(\frac{1}{h} - \beta)$. If we now choose $h = 1/\beta d$, we can check that $\kappa \leq \exp(4/d)$ for $d \geq 2$. By Proposition 7.1.2, if we have access to the minimizer of V_y (which is equivalent to being able to compute the proximal oracle for hV), we can construct an upper envelope for which the average number of iterations of rejection sampling is bounded by $\kappa^{d/2} \leq \exp(2) \leq 8$. We summarize this discussion as follows.

Implementing the RGO via Rejection Sampling: To sample from $\pi^{X|Y}(\cdot \mid y)$, where V is β -smooth, we proceed via the following steps.

- 1 Compute the minimizer x_y^* of V_y defined via $V_y(x) := V(x) + \frac{1}{2h} \|y - x\|^2$, and compute the minimum value V_y^* . This can be done exactly if we assume access to a proximal oracle for V (which is a natural assumption when designing a proximal algorithm for sampling); otherwise, if $h < 1/\beta$, then this is a strongly convex optimization problem and can be implemented using standard algorithms.
- 2 Let $\tilde{\pi}^{X|Y}(\cdot \mid y) := \exp\{-(V_y - V_y^*)\}$ and $\tilde{\mu}_y := \exp(-\frac{1/h - \beta}{2} \|\cdot - x_y^*\|^2)$. Use rejection sampling to sample from $\pi^{X|Y}$ with the envelope $\tilde{\mu}_y$.

Each iteration of rejection sampling requires one call to an evaluation oracle for V (in order to compute the acceptance probability). We summarize the guarantees for this implementation of the RGO in the following theorem.

Theorem 8.6.1 (Naïve rejection sampling implementation). *Assume that V is β -smooth. Then, if $h \leq 1/\beta d$, rejection sampling implements the RGO for π^X exactly using one computation to a proximal oracle for V and $O(1)$ expected calls to an evaluation oracle for V .*

Although the rejection sampling strategy is naïve, it yields surprisingly strong results with a simple analysis. If we assume that π^X satisfies an LSI with constant $1/\alpha$, then by Theorem 8.5.2 yields an iteration complexity of roughly $\tilde{O}(\frac{1}{\alpha h} \log(1/\varepsilon)) = \tilde{O}(\kappa d \log(1/\varepsilon))$, where $\kappa := \beta/\alpha$ and each iteration requires $O(1)$ queries in expectation. This result already matches the guarantees for MALA under a cold start (Theorem 7.3.4), and under a weaker assumption (LSI rather than strong log-concavity)! It is also straightforward to obtain guarantees in Rényi divergence (Exercise 8.6), which are stronger than Theorem 6.1.2 (by a factor of κ , and because the result only has a logarithmic dependence on $1/\varepsilon$). Finally, under a Poincaré inequality, Theorem 8.5.1 yields an iteration complexity of roughly $\tilde{O}(\kappa d^2 \log(1/\varepsilon))$ (assuming that $\chi^2(\mu_0 \parallel \pi) \leq \exp(\tilde{O}(d))$), which moreover can be shown to hold in any Rényi divergence, and which can be compared to the $\tilde{O}(\kappa^2 d^3/\varepsilon^2)$ rate obtained via LMC (Theorem 6.2.9) or the $\tilde{O}(\kappa^{3/2} d^2/\varepsilon)$ rate obtained via ULMC (Corollary 6.3.2) in this setting.

Implementation via high-accuracy samplers.

Next, we consider the RGO with step size $h = 1/2\beta$, so that the RGO is a strongly log-concave and log-smooth distribution with condition number κ . In this regime, the iteration complexity of the proximal sampler is roughly $\tilde{O}(\kappa \log(1/\varepsilon))$ under an LSI, or $\tilde{O}(\kappa d \log(1/\varepsilon))$ under a PI, but we also have to take into account the per-iteration complexity of implementing the RGO.

Here, the implementation of the RGO will be handled by one of the high-accuracy samplers from Chapter 7. However, in order to complete the analysis, we also need to analyze how the error propagates due to the imperfect implementation. Here, we state a very simple version of the error analysis, which is proven via coupling (Exercise 8.8).

Lemma 8.6.2. *Let μ_N^X denote the law of the N -th iterate of the proximal sampler with perfect implementation of the RGO. Suppose that instead, in each step of the proximal sampler, we use a sample from a distribution which is δ -close to the RGO in total variation distance; let $\hat{\mu}_N^X$ denote the law of the N -th iterate of the proximal sampler with imperfect implementation of the RGO. Then, it holds that*

$$\|\hat{\mu}_N^X - \mu_N^X\|_{\text{TV}} \leq N\delta.$$

In any case, we only need to run the proximal sampler for a number of iterations N which is polynomial in the various problem parameters, which means δ can be taken to be inverse polynomially small. But the high-accuracy guarantees from Chapter 7 can sample from the RGO to accuracy δ with a number of iterations that only scales logarithmically with $1/\delta$, which ultimately contributes a negligible overhead to the analysis if we ignore logarithmic factors. Therefore, the total complexity of the proximal sampler is roughly given by its iteration complexity, multiplied by the dimension dependence of the high-accuracy sampler (note that the condition number of the RGO is $O(1)$ so it does not play a role here). In summary, the use of a high-accuracy sampler allows us to safely neglect the errors arising from inexact implementation, but one could also consider other implementations with a more careful error analysis.

In particular, let us consider the state-of-the-art high-accuracy sampler—MALA initialized with

a **ULMC** warm start—from [Corollary 7.3.6](#). That result required knowing the mode of the distribution, which for the RGO amounts to a call to a proximal oracle for V . The complexity of implementing the RGO is then $\tilde{O}(\sqrt{d} \text{polylog}(1/\varepsilon))$, which readily yields the following corollaries. We focus solely on total variation guarantees for simplicity; see [Altschuler and Chewi \(2024a\)](#) for more detailed results.

Corollary 8.6.3 (High-accuracy sampling via proximal sampler). *Let $\pi^X \propto \exp(-V)$, where V is β -smooth. Take $h = 1/2\beta$ and assume we have an evaluation and proximal oracle for V . Let $\hat{\mu}_N^X$ denote the law of the N -th iterate of the proximal sampler, in which the RGO is implemented via [Corollary 7.3.6](#).*

1 ([Theorem 8.4.1](#)) *If V is convex, we obtain the guarantee $\|\hat{\mu}_N^X - \pi^X\|_{\text{TV}} \leq \varepsilon$ using*

$$\tilde{O}\left(\frac{\beta\sqrt{d} W_2^2(\hat{\mu}_0^X, \pi^X)}{\varepsilon^2}\right) \quad \text{queries}.$$

2 ([Theorem 8.5.1](#)) *If π^X satisfies a Poincaré inequality with constant $1/\alpha$, and we write $\kappa := \beta/\alpha$ for the condition number, we obtain $\|\hat{\mu}_N^X - \pi^X\|_{\text{TV}} \leq \varepsilon$ using*

$$O(\kappa\sqrt{d} \{\log \chi^2(\hat{\mu}_0^X \parallel \pi) + \text{polylog}(1/\varepsilon)\}) \quad \text{queries}.$$

3 ([Theorem 8.5.2](#)) *If π^X satisfies a log-Sobolev inequality with constant $1/\alpha$, and we write $\kappa := \beta/\alpha$ for the condition number, we obtain $\|\hat{\mu}_N^X - \pi^X\|_{\text{TV}} \leq \varepsilon$ using*

$$\tilde{O}\left(\kappa\sqrt{d} \text{polylog} \frac{\text{KL}(\hat{\mu}_0^X \parallel \pi)}{\varepsilon^2}\right) \quad \text{queries}.$$

These guarantees are state-of-the-art for each of their corresponding settings.

Direct implementation for composite log-concave sampling.

In an illuminating work, [Fan et al. \(2023\)](#) developed an approximate rejection sampling scheme, inspired by the one in [Gopi et al. \(2022\)](#), which succeeds at implementing the RGO with nearly constant query complexity, for step size $h \asymp 1/\beta\sqrt{d}$. By combining this with the convergence results for the proximal sampler established in this chapter, this yields another approach to obtaining the guarantees of [Corollary 8.6.3](#) while completely bypassing the analysis of MALA!

In this section, following [Liu and Chewi \(2026\)](#), we present a generalization to the composite setting where $V = f + g$, and g is convex but potentially non-smooth.

Optimization Box 8.6.4. Consider a composite objective $V := f + g$, where f is α_f -convex and β -smooth, and g is α_g -convex, non-smooth, but “simple” so that the proximal oracle for g can be evaluated in closed form. The canonical example is when $g = \|\cdot\|_1$ is the ℓ_1 -norm, in which case this composite objective encompasses computation of the LASSO estimator for sparse regression. In this case, the **proximal gradient** algorithm consists of GD steps for f and PPM steps for g : $x_{n+1} = \text{prox}_{hg}(x_n - h \nabla f(x_n))$.

The proximal gradient algorithm can also be written as

$$x_{n+1} = \arg \min_{x \in \mathbb{R}^d} \psi_{x_n}(x) := \arg \min_{x \in \mathbb{R}^d} \left\{ f(x_n) + \langle \nabla f(x_n), x - x_n \rangle + g(x) + \frac{1}{2h} \|x - x_n\|^2 \right\}.$$

Since ψ_{x_n} is $(\alpha_g + 1/h)$ -convex, it grows quadratically around its minimizer: for any $y \in \mathbb{R}^d$, $\psi_{x_n}(y) \geq \psi_{x_n}(x_{n+1}) + \frac{\alpha_g + 1/h}{2} \|y - x_{n+1}\|^2$. Also, from our assumptions on f and g ,

$$\psi_{x_n}(y) \leq V(y) + \frac{1/h - \alpha_f}{2} \|y - x_n\|^2, \quad \psi_{x_n}(x_{n+1}) \geq V(x_{n+1})$$

for $h \leq 1/\beta$. Combining these inequalities yields the one-step bound

$$(1 + \alpha_g h) \|y - x_{n+1}\|^2 \leq (1 - \alpha_f h) \|y - x_n\|^2 - 2h \{V(x_{n+1}) - V(y)\}, \quad (8.6.5)$$

which generalizes (4.3.3). In particular, the proximal gradient algorithm converges at the same rate as GD for convex *smooth* optimization, despite the presence of the non-smooth term.

As discussed in [Optimization Box 8.0.1](#), although non-smoothness typically slows down the convergence rates for optimization, this is mitigated if we have access to prox_{hg} via the proximal gradient method. In the context of sampling, the analogous assumption is to suppose that we have access to the RGO for g , denoted $\text{RGO}_{g,h,y}$:

$$\text{RGO}_{g,h,y}(x) \propto \exp\left(-g(x) - \frac{1}{2h} \|y - x\|^2\right).$$

Note that when g is the convex indicator for a convex set $\mathcal{C}$, i.e., $g = 0$ on $\mathcal{C}$ and $g = +\infty$ outside $\mathcal{C}$, this corresponds to sampling from a Gaussian conditioned to lie in $\mathcal{C}$. Other examples, typically separable ones (e.g., including $g = \lambda \|\cdot\|_1$), are discussed in [Liu and Chewi \(2026\)](#).

Our goal is to prove the following theorem.

Theorem 8.6.6 (High-accuracy composite sampling). *Let $\pi^X \propto \exp(-V) = \exp(-f - g)$, where we assume:*

- f is α_f -convex and β -smooth.
- g is α_g -convex, $\alpha_g \geq 0$.

We assume that f and g are both minimized at 0, $\alpha := \alpha_f + \alpha_g > 0$, and we write $\kappa := 1 \vee \beta/\alpha$. Then, there is an implementation of the proximal sampler which yields a sample ε -close to π^X in total variation distance using

$$N = \tilde{O}(\kappa \sqrt{d} \text{polylog}(1/\varepsilon)) \quad \text{queries to } f, \nabla f, \text{ and } \text{RGO}_g.$$

Note that by taking $g = 0$, we recover the complexity result of [Corollary 8.6.3](#) in the strongly log-concave case. Extensions to other settings are presented in [Exercise 8.12](#).

The assumption that f and g are both minimized at 0 is without loss of generality. Indeed, if this does not hold, then compute $x_\star := \arg \min(f + g)$ (e.g., via proximal gradient). The optimality condition says that $-\nabla f(x_\star) \in \partial g(x_\star)$, so by replacing f with $f - \langle \nabla f(x_\star), \cdot \rangle$ and g with $g + \langle \nabla f(x_\star), \cdot \rangle$, we ensure that f and g are both minimized at $x_\star$, and then we can translate to make $x_\star$ the origin.

The idea is to apply the proximal sampler with target π , which requires implementing

$$\text{RGO}_{f+g,h,y}(x) \propto \exp\left(-f(x) - g(x) - \frac{1}{2h} \|y - x\|^2\right).$$

To do so, consider approximating this distribution via

$$\text{RGO}_{g,h,y-h\nabla f(y)}(x) \propto \exp\left(-\langle \nabla f(y), x - y \rangle - g(x) - \frac{1}{2h} \|y - x\|^2\right).$$

Thus, we linearize the smooth term, similarly to the proximal gradient method. This approximation incurs some error, so we correct it via a filter: namely, we use $\text{RGO}_{g,h,y-h\nabla f(y)}$ as a proposal in the independent Metropolis–Hastings algorithm with target $\text{RGO}_{f+g,h,y}$ (see [Exercise 7.11](#)).

We initialize the entire algorithm using the following lemma.

Lemma 8.6.7 (Composite initialization). *Under the assumptions of [Theorem 8.6.6](#), the initialization $\mu_0^X := \text{RGO}_{g,1/\beta,0}$ satisfies*

$$\mathcal{R}_\infty(\mu_0^X \parallel \pi^X) \leq \frac{d}{2} \log \kappa.$$

See [Exercise 8.10](#) or [Lee et al. \(2021b, Proposition 61\)](#).

Proof of [Theorem 8.6.6](#) By [Theorem 8.5.2](#), using the initialization in [Lemma 8.6.7](#), we see that the ideal proximal sampler (i.e., with perfect RGO implementation) converges in at most $N_{\text{outer}} = \tilde{O}((1/\alpha h) \log(1/\varepsilon))$ iterations. Let $h \ll 1/\beta$.

Implementation of $\text{RGO}_{f+g,h,y}$ for fixed y . First, by [Exercise 7.11](#), the independent Metropolis–Hastings implementation yields a δ -accurate sample in TV distance from $\text{RGO}_{f+g,h,y}$ in at most $N_{\text{inner}} = O(\log(1/\delta))$ iterations, provided that we can bound the Rényi divergence between the proposal and the target in both directions. Toward this end, we apply [Exercise 6.3](#), noting that $\text{RGO}_{f+g,h,y}$, $\text{RGO}_{g,h,y-h\nabla f(y)}$ both satisfy an LSI with constant at most $O(h)$. Choosing $\lambda \asymp hq^2$, for $X \sim \text{RGO}_{g,h,y-h\nabla f(y)}$,

$$\begin{aligned} \mathcal{R}_q(\text{RGO}_{g,h,y-h\nabla f(y)} \parallel \text{RGO}_{f+g,h,y}) &\leq \log \mathbb{E} \exp O\left(hq^2 \left\| \nabla \log \frac{\text{RGO}_{g,h,y-h\nabla f(y)}(X)}{\text{RGO}_{f+g,h,y}} \right\|^2\right) \\ &= \log \mathbb{E} \exp O\left(hq^2 \|\nabla f(X) - \nabla f(y)\|^2\right) \\ &\leq \log \mathbb{E} \exp O(\beta^2 hq^2 \|X - y\|^2). \end{aligned}$$

Let $y^+ := \text{prox}_{hg}(y - h\nabla f(y))$ denote one step of the proximal gradient method starting from y , and note that y^+ is the mode of $\text{RGO}_{g,h,y-h\nabla f(y)}$. By [Lemma 4.0.1](#) and sub-Gaussian concentration ([Theorem 2.4.3](#)),

$$\begin{aligned} \log \mathbb{E} \exp O(\beta^2 hq^2 \|X - y\|^2) &\leq \beta^2 hq^2 \|y^+ - y\|^2 + \log \mathbb{E} \exp O(\beta^2 hq^2 \|X - y^+\|^2) \\ &\leq \beta^2 hq^2 \|y^+ - y\|^2 + \beta^2 hq^2 \mathbb{E}[\|X - y^+\|^2] + \log \mathbb{E} \exp O(\beta^2 hq^2 \{\|X - y^+\| - \mathbb{E}_\pi\|X - y^+\|\}^2) \\ &\leq \beta^2 hq^2 \|y^+ - y\|^2 + \beta^2 dh^2 q^2 + \beta^2 h^2 q^2 \leq \beta^2 hq^2 \|y^+ - y\|^2 + \beta^2 dh^2 q^2 \end{aligned}$$

provided $h \ll 1/\beta q$. The reverse bound on $\mathcal{R}_q(\text{RGO}_{f+g,h,y} \parallel \text{RGO}_{g,h,y-h\nabla f(y)})$ is the same and follows from a similar argument.

In the context of [Exercise 7.11](#), we can identify $C_R \asymp \beta^2 h \|y^+ - y\|^2 + \beta^2 d h^2$ and $\gamma = 2$. We conclude that independent Metropolis–Hastings succeeds at implementing $\text{RGO}_{f+g,h,y}$, up to error δ in TV distance, in $O(\log(1/\delta))$ iterations, provided that

$$h \ll \frac{1}{\beta \log(1/\delta)} \quad \text{and} \quad \beta^2 h \|y^+ - y\|^2 + \beta^2 d h^2 \ll \frac{1}{\log^3(1/\delta)}. \quad (8.6.8)$$

Control of $\|y^+ - y\|$ along the iterates. Next, we provide a high-probability bound on the quantity $\|y^+ - y\|$ along the iterates of the proximal sampler. By [\(8.6.5\)](#),

$$\|y^+ - y\|^2 \leq \frac{2h}{1 - \alpha_f h} \{V(y) - V(y^+)\} \leq 4h \{V(y) - V_\star\},$$

where $h \leq 1/2\beta \leq 1/2|\alpha_f|$ and $V_\star := \inf V$. We apply [Lemma 8.6.10](#) below, which implies that for $X \sim \pi^X$, with probability at least $1 - \delta$,

$$V(X) - V_\star \lesssim d + \log(1/\delta).$$

Then, we can write $Y \sim \pi^Y$ as $Y = X + \sqrt{h}\xi$, where $\xi \sim \text{normal}(0, I_d)$. Since prox_{hg} is 1-Lipschitz ([Exercise 8.3](#)) and $\text{id} - h\nabla f$ is 2-Lipschitz under our assumptions,

$$\begin{aligned} \|Y^+ - Y\|^2 &\lesssim \|Y^+ - X^+\|^2 + \|X^+ - X\|^2 + \|Y - X\|^2 \lesssim \|Y - X\|^2 + \|X^+ - X\|^2 = h \|\xi\|^2 + \|X^+ - X\|^2 \\ &\lesssim (d + \log(1/\delta)) h, \end{aligned}$$

again with probability at least $1 - \delta$. This establishes concentration for the stationary distribution.

Next, let $(\mu_n^Y)_{n \in \mathbb{N}}$ denote the Y -marginals of the ideal proximal sampler. By the data-processing inequality ([Theorem 1.5.6](#)) and [Lemma 8.6.7](#), we have $\mathcal{R}_\infty(\mu_n^Y \parallel \pi^Y) \leq \frac{d}{2} \log \kappa$ for all n . Hence, by change of measure ([Lemma 6.2.8](#)), we deduce that for $Y_n \sim \mu_n^Y$, with probability at least $1 - \delta$,

$$\|Y_n^+ - Y_n\|^2 \lesssim \{d(1 + \log \kappa) + \log(1/\delta)\} h. \quad (8.6.9)$$

Choose $\delta = \varepsilon/N_{\text{outer}}$ so that by a union bound, $\|Y_n^+ - Y_n\|$ satisfies the desired bound for all n with probability at least $1 - \varepsilon$. Substituting this into [\(8.6.8\)](#), we see that if $h = \tilde{\Theta}(1/\beta \{\sqrt{d} \log^{3/2}(1/\varepsilon) + \log^2(1/\varepsilon)\})$, then the independent Metropolis–Hastings subroutine will succeed at implementing RGO_{f+g,h,Y_n} for all n , using $O(\log(1/\delta))$ steps each, with probability at least $1 - \varepsilon$.

Completing the proof. The remaining step to make this proof rigorous is a coupling argument similar to [Lemma 8.6.2](#). The core logic is that we iteratively build a coupling between the ideal and implemented proximal samplers, such that the corresponding iterates are equal at each step with high probability. At each step, we accumulate a failure probability of size $O(\delta)$ due to the event that Y_n fails to satisfy the concentration bound [\(8.6.8\)](#), and from the TV error incurred by independent Metropolis–Hastings. We leave the details to the reader (or see [Liu and Chewi, 2026](#)).

The total query complexity is

$$N_{\text{inner}} \times N_{\text{outer}} = \tilde{O}(\kappa \{\sqrt{d} \log^{7/2}(1/\varepsilon) + \log^4(1/\varepsilon)\}). \quad \square$$

In the proof, we used the following concentration bound; see [Exercise 8.11](#).

Lemma 8.6.10 (Concentration of the potential). *Let V be convex and $\pi \propto \exp(-V)$. Then, for any $0 < \lambda < 1$, $\mathbb{E}_\pi \exp(\lambda(V - V_\star)) \leq 1/(1 - \lambda)^d$, where $V_\star := \inf V$.*

In particular, for any $0 < \delta < 1$, with probability at least $1 - \delta$ under π ,

$$V - V_\star \leq d + \sqrt{2d \log(1/\delta)} + \log(1/\delta).$$

Sampling from general log-concave densities over convex sets.

There has also been interest in using the proximal sampler to sample from the uniform distribution, or more generally a log-concave density, over a convex set $\mathcal{C}$. Here, we only assume access to zeroth-order evaluation, that is, queries to a membership oracle for $\mathcal{C}$ and to evaluations of the density. This line of work began with [Kook et al. \(2024\)](#), continued in [Kook \(2025\)](#); [Kook and Vempala \(2025a,c\)](#); [Kook and Zhang \(2025\)](#); [Kook and Vempala \(2026\)](#), and was extended to non-convex bodies in [Vempala and Wibisono \(2026\)](#). Note that for $\pi^X = \text{uniform}(\mathcal{C})$, implementation of the RGO boils down to sampling from a Gaussian truncated to $\mathcal{C}$, which can be quite challenging. In particular, bounding the rejection probability when using rejection sampling with a simple proposal involves geometric considerations of the set in question. Nevertheless, this strategy has proven to be extremely successful, leading to the following state-of-the-art guarantee.

Theorem 8.6.11 ([Kook \(2025\)](#)). *Let $\mathcal{C}$ be a convex body containing $B(0, 1)$, and let $\pi := \text{uniform}(\mathcal{C})$. For Σ the covariance matrix of π , let $\Lambda := \|\Sigma\|_{\text{op}}$ and $R^2 := \text{tr } \Sigma$. Then, there is an algorithm such that for any $q \geq 2$, $\varepsilon > 0$, and any initialization μ_0 , the algorithm outputs a sample from $\hat{\mu}$ satisfying $\mathcal{R}_q(\hat{\mu} \parallel \pi) \leq \varepsilon^2$ using*

$$\tilde{O}\left(qd^2 R^{3/2} \Lambda^{1/4} + qd^2 \Lambda \log \frac{1}{\varepsilon}\right) \quad \text{membership queries in expectation.}$$

The first term represents the cost of producing a warm start, and the second term is the iteration complexity given a warm start. The result can also be extended to general log-concave densities, see [Kook \(2025\)](#) for details.

Bibliographical Notes

The reader is encouraged to read the original paper [Lee et al. \(2021c\)](#) on the proximal sampler, which contains applications to sampling from composite densities $\pi \propto \exp(-(f + g))$, where f is well-conditioned and g admits an implementable RGO, as well as to sampling from log-concave finite sums $\pi \propto \exp(-F)$ where $F := n^{-1} \sum_{i=1}^n f_i$ is well-conditioned and the complexity is measured via the number of oracle calls to the individual functions $(f_i)_{i \in [n]}$. The proximal sampler has also been used to sample from weakly smooth and non-smooth potentials ([Liang and Chen, 2022a,b](#)), and it has been applied to the problem of differentially private convex optimization ([Gopi et al., 2022](#)). A variant of the proximal sampler tailored to handle non-Euclidean norms was proposed in [Gopi et al. \(2023\)](#), although it is currently less well-understood than its Euclidean counterpart.

The interpretation of the proximal sampler in Section 8.1 arose out of conversations with Adil Salim, and was inspired by the interpretation of LMC in [Wibisono \(2018\)](#). The simultaneous flow calculation ([Theorem 8.3.1](#)) has appeared in many works, but the version that appears here, emphasizing the Markov semigroup perspective and the role of the carré du champ, was anticipated by Ramon van

Handel. The main results in this chapter, on the convergence of the proximal sampler, are from [Lee et al. \(2021b\)](#); [Chen et al. \(2022b\)](#).

As we have seen, the proximal sampler can be interpreted as the time reversal of the heat flow, which itself was studied in [Klartag and Putterman \(2023\)](#); this paper also introduced the adjoint of the heat semigroup ([Exercise 8.2](#)) and connected it with the Föllmer process ([Föllmer, 1985](#)) and Eldan’s stochastic localization ([Eldan, 2013](#)). This was further explored in [Chen and Eldan \(2022\)](#), which interpreted stochastic localization as the “localization scheme” (a concept introduced in the paper) associated with the proximal sampler and used this to give another proof of convergence in KL divergence under strong log-concavity. The adjoint of the heat semigroup was also important for recent progress on the KLS conjecture ([Klartag and Lehec, 2022](#)). See the Bibliographical Notes to Chapter 3 for further references.

The optimization results in [Exercise 8.3](#) obtained in analogy with the proximal sampler are given in [Chen et al. \(2022b\)](#).

Exercises

Introduction to the Proximal Sampler

▷ **Exercise 8.1** (RGO as a proximal operator on the Wasserstein space)

Given a functional $\mathcal{F} : \mathcal{P}_2(\mathbb{R}^d) \rightarrow \mathbb{R} \cup \{\infty\}$, the proximal operator for $\mathcal{F}$ on the Wasserstein space is defined via

$$\text{prox}_{\mathcal{F}}(\mu) := \arg \min_{\mu' \in \mathcal{P}_2(\mathbb{R}^d)} \left\{ \mathcal{F}(\mu') + \frac{1}{2} W_2^2(\mu, \mu') \right\}.$$

The proximal operator was used in the seminal work [Jordan et al. \(1998\)](#) in order to rigorously make sense of gradient flows on the Wasserstein space. Prove that the RGO satisfies

$$\pi^{X|Y=y} = \text{prox}_{h \text{KL}(\cdot \| \pi)}(\delta_y).$$

Hence, the assumption that we can implement the RGO is the same as assuming that we can evaluate the proximal operator for the KL divergence on any Dirac measure.

▷ **Exercise 8.2** (RGO as an adjoint)

Let Q_h denote the Markov kernel corresponding to the heat flow for time h , i.e., $Q_h f(x) := \mathbb{E} f(x + \sqrt{h} \xi)$, where $\xi \sim \text{normal}(0, I_d)$. Show that Q_h defines a bounded linear map $L^2(\pi^Y) \rightarrow L^2(\pi^X)$ (in fact, with operator norm equal to 1).

Show that the adjoint Q_h^* of this linear map is precisely the RGO.

Convergence under Strong Log-Concavity

▷ **Exercise 8.3** (Comparison with optimization results)

This exercise compares the results for the proximal sampler with the proximal point method in optimization.

1 Suppose that V is α -strongly convex. Prove that prox_{hV} is $\frac{1}{1+\alpha h}$ -Lipschitz.

Hint: Show that $\text{prox}_{hV} = (\text{id} + h\nabla V)^{-1}$. Argue via convex duality by considering the convex conjugate $(\frac{\|\cdot\|^2}{2} + hV)^*$.

2 Suppose that V is α -strongly convex. Translate both of the proofs in Section 8.2 to Euclidean optimization.

3 Suppose that V satisfies the gradient domination condition

$$\|\nabla V(x)\|^2 \geq 2\alpha \{V(x) - \inf V\}, \quad \text{for all } x \in \mathbb{R}^d.$$

Also, let $x' := \text{prox}_{hV}(x)$. Inspired by [Theorem 8.5.2](#), we can ask whether or not it holds that

$$V(x') - \inf V \leq \frac{1}{(1 + \alpha h)^2} \{V(x) - \inf V\}.$$

Prove that this is indeed the case.

Hint: Define $V_{t,x}(z) := V(z) + \frac{1}{2t} \|z - x\|^2$ and let $x_t := \arg \min V_{t,x}$; then, differentiate $t \mapsto V_{t,x}(x_t)$.

▷ [Exercise 8.4](#) (First coupling lemma)

Prove Lemma 8.2.2.

Simultaneous Flow and Time Reversal

▷ [Exercise 8.5](#) (Convergence under LSI)

Verify that the result under LSI ([Theorem 8.5.2](#)) holds.

▷ [Exercise 8.6](#) (Convergence in Rényi divergence)

Following [Theorem 2.2.25](#), extend [Theorem 8.5.1](#) and [Theorem 8.5.2](#) to provide Rényi divergence guarantees.

▷ [Exercise 8.7](#) (Gaussian case)

In this exercise, we consider the Gaussian case for intuition.

- 1 Suppose that $\pi^X = \text{normal}(0, \Sigma)$ and $\mu_0^X = \text{normal}(m_0, \Sigma_0)$. Show that the iterates of the proximal sampler all have Gaussian distributions, and explicitly compute the means and variances.
- 2 In particular, show that $\mu_1^X = \text{normal}(m_1, \Sigma_1)$, where the mean satisfies $m_1 = \text{prox}_{hV}(m_0)$ and $V(x) := \frac{1}{2} \langle x, \Sigma^{-1} x \rangle$. In other words, the mean of the iterate of the proximal sampler evolves according to the proximal point method. Use this to show that the contraction factors in [Theorem 8.2.1](#) and [Theorem 8.5.2](#) are sharp.

Implementations of the RGO and Applications

▷ [Exercise 8.8](#) (Second coupling lemma)

Prove Lemma 8.6.2.

▷ [Exercise 8.9](#) (MHMC implementation)

Consider using [Theorem 7.3.7](#) to implement the RGO. Show that in this way, the complexity bound of [Theorem 7.3.7](#) can be improved to $\tilde{O}((\kappa + \bar{\kappa}^{2/3}) d^{1/4} \text{polylog}(1/\varepsilon))$. Extend this result to other settings (LSI, PI, weakly log-concave).

▷ [Exercise 8.10](#) (Composite initialization)

Prove Lemma 8.6.7.

▷ [Exercise 8.11](#) (Concentration of the potential)

Prove Lemma 8.6.10.

▷ **Exercise 8.12** (Composite sampling in more general settings)

Here, we extend the result of Theorem 8.6.6 to other settings.

- 1 First, suppose that f is β -smooth but not necessarily convex; however, we keep the assumption that g is α_g -convex. Derive convergence bounds in the three cases: π^X is log-concave; π^X satisfies a PI; π^X satisfies an LSI.

Hint: Use Theorem 8.6.6 to implement $\text{RGO}_{f+g,h,y}$.

- 2 Next, instead of assuming that f is β -smooth, assume that f is L -Lipschitz. Adapt the argument of the implementation of the RGO to show that a step size of $h \asymp 1/L^2$ suffices, up to logarithmic factors. Then, derive convergence bounds in the same three settings as above.

Hint: The LSI for $\text{RGO}_{f+g,h,y}$ comes from Proposition 2.3.4 in this case.

Lower Bounds for Sampling

In order to determine if our sampling guarantees are optimal, we need to pair them with matching lower bounds. However, the problem of establishing query complexity lower bounds for sampling has been challenging and the work on this topic is nascent. Here, we will give an overview of the current progress in this direction.

Optimization Box 9.0.1. In convex optimization, the situation is better understood. For example, in the high-dimensional regime, the complexity of finding an ε -minimizer of a function $V : \mathbb{R}^d \rightarrow \mathbb{R}$ with $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and minimizer $x_\star$ satisfying $\|x_\star\| \leq R$ is $\Omega(\sqrt{\kappa} \log(R^2/\alpha\varepsilon))$, where $\kappa := \beta/\alpha$ (Nemirovsky and Yudin, 1983). This matches the complexity of accelerated gradient descent (Optimization Box 5.3.2), which is therefore optimal for this function class. In this case, the lower bound was actually a motivating force which led to the discovery of the optimal algorithm.

Subsequently, lower bounds for convex optimization have been established for a host of function classes and oracle models.

9.1 A Brief Discussion of Query Complexity

Before turning to the results, recall the discussion and interpretation of *query complexity* (or *oracle complexity*) from the Preface, although here we present the model in more detail. Let Π be a family of probability measures over $\mathbb{R}^d$ with continuously differentiable densities. A *first-order oracle* for $\pi \in \Pi$ with density $\pi \propto \exp(-V)$ is the mapping $\mathcal{O}_\pi : \mathbb{R}^d \rightarrow \mathbb{R} \times \mathbb{R}^d$ given by $x \mapsto (V(x) - V(0), \nabla V(x))$. A *first-order sampling algorithm* interacts with the oracle $\mathcal{O}_\pi$ by querying N points $x_1, \dots, x_N$, where each query point x_j depends on the previous query points and query values $\{x_i, \mathcal{O}_\pi(x_i)\}_{i \in [j-1]}$ as well as an independent source of randomness, and outputs a “sample” x_{out} . If x_{out} is the random output of the sampling algorithm when run on oracle $\mathcal{O}_\pi$, we say that the sampling algorithm computes the measure $\text{law}(x_{\text{out}})$.

The *first-order complexity* $\mathcal{C}(\Pi, d, \varepsilon)$ of the class Π is the minimum number N for which there exists a sampling algorithm with the following property: for any $\pi \in \Pi$, when given N queries to

oracle $\mathcal{O}_\pi$, the sampling algorithm computes $\hat{\pi}$ with $d(\hat{\pi}, \pi) \leq \varepsilon$. Here, d can be taken to be total variation distance or the Wasserstein distance, or even an asymmetric measure such as $\sqrt{\text{KL}}$. Our goal is to characterize $\mathcal{C}(\Pi, d, \varepsilon)$ up to universal constants for important classes Π . In particular, we will be interested in the class of strongly log-concave and log-smooth distributions over $\mathbb{R}^d$ with known mode at 0, since we view this class as a fundamental benchmark with which to test our understanding of acceleration, discretization, etc.

One can also consider various extensions, e.g., by changing the class of measures Π or by changing the oracle to allow for different information structures (zeroth-order oracles, second-order oracles, etc.), finite sum structures, stochasticity, and so on. It is important to eventually understand the complexities for all of these settings—and indeed they have been carefully mapped out in the sister field of optimization—but we emphasize that for sampling, even the canonical case of well-conditioned log-concave measures with an exact oracle is still not fully understood.

We also recall that query complexity is not computational complexity: our model allows the algorithm to perform an unbounded amount of computation in between queries to the oracle. This is precisely the strength of this model: due to its information-theoretic nature, it is possible to establish unconditional lower bounds. In practice, the two are closely connected. If the only access one has to the target density is through the oracle (as opposed to other forms of access, such as integrals of certain test functions, etc.), then query complexity is a lower bound for the computational complexity. On the other hand, reasonable algorithms with small query complexity tend to also have small computational complexity (indeed this is true for all of the algorithms we have encountered in this book thus far), and so we treat query complexity as a useful proxy for computation time.

Finally, we briefly discuss differences with the setting of optimization. Optimization lower bounds are often first established for simple models of deterministic algorithms, e.g., algorithms whose iterates always lie in the span of the queried gradients (Nesterov, 2018). These lower bounds can sometimes be extended to randomized algorithms with additional tricks. In sampling, we do not have this luxury, since randomness is an integral part of any sampler. Moreover, whereas lower bound constructions for optimization must hide the location of the minimizer from the algorithm, lower bound constructions for sampling must hide *the bulk of the distribution's mass* from the algorithm, which may explain the difficulty of the constructions.

9.2 Query Complexity in One Dimension

The first query complexity result for well-conditioned log-concave sampling is from Chewi et al. (2022), which established a sharp result in dimension one. In this section, we present this result in detail because it is simple enough to do so.

Define the class

$$\Pi_{\kappa,1} := \{\pi \in \mathcal{P}_{\text{ac}}(\mathbb{R}) \mid \pi \propto \exp(-V), \ 1 \leq V'' \leq \kappa, \ V'(0) = 0\}.$$

In our definition of $\Pi_{\kappa,1}$, we have enforced the requirement $V'(0) = 0$ in order to cleanly separate out the complexity of optimization (finding the minimizer of V) from the intrinsic complexity of sampling. Note that some restriction of this nature is needed, since without *any* prior knowledge of the mode it is impossible to find it.

Theorem 9.2.1 (Chewi et al. (2022)). *The query complexity of outputting a sample which is $\frac{1}{64}$ close in total variation distance to the target π , uniformly over the choice of $\pi \in \Pi_{\kappa,1}$, is $\Theta(\log \log \kappa)$. In other words, $\mathcal{C}(\Pi_{\kappa,1}, \|\cdot\|_{\text{TV}}, \frac{1}{64}) \asymp \log \log \kappa$.*

Lower bound.

The lower bound will proceed in two stages.

- 1 First, we reduce the sampling problem to a statistical testing problem. Namely, we will construct a family $\pi_1, \dots, \pi_m \in \Pi_{\kappa,1}$, and suppose that $i \sim \text{uniform}([m])$ is drawn randomly. The statistical testing problem is defined as follows: given query access to π_i (through the oracle), guess the value of i . We will show that an algorithm to sample from π_i can be used to solve the statistical testing problem; thus, “sampling is harder than testing”.
- 2 Next, we will prove a lower bound on the number of queries required to solve the statistical testing problem: “testing is hard”. This relies on standard information-theoretic techniques for proving minimax lower bounds for statistical problems. The main difference between this problem and the usual statistical setting is that rather than having i.i.d. samples from some data distribution, we instead have query access and the algorithm is allowed to be adaptive.

Combining the two steps then yields our query complexity lower bound for sampling. We begin with the construction of $\pi_1, \dots, \pi_m$, which is slightly tricky.

Let m be the largest integer such that $\exp(-\frac{2^{2m-2}}{2\kappa}) \geq \frac{1}{2}$ (and note that $m = \Theta(\log \kappa)$). We define a family $\{V_i\}_{i \in [m]}$ of 1-strongly convex and κ -smooth potentials as follows. We require that $V_i(0) = V'_i(0) = 0$ and that V_i be an even function. The second derivative is

$$V''_i(x) := \begin{cases} \kappa, & \sqrt{\kappa}|x| \in [2^{i-1}, 2^i) \cup \bigcup_{j=i+1}^{m+1} [2^j, \frac{5}{4}2^j), \\ 1, & \text{otherwise.} \end{cases}$$

Although the construction seems complicated, the basic idea is to make V''_i oscillate between its minimum and maximum allowable values 1 and κ ; see Figure 9.1 for a visual.

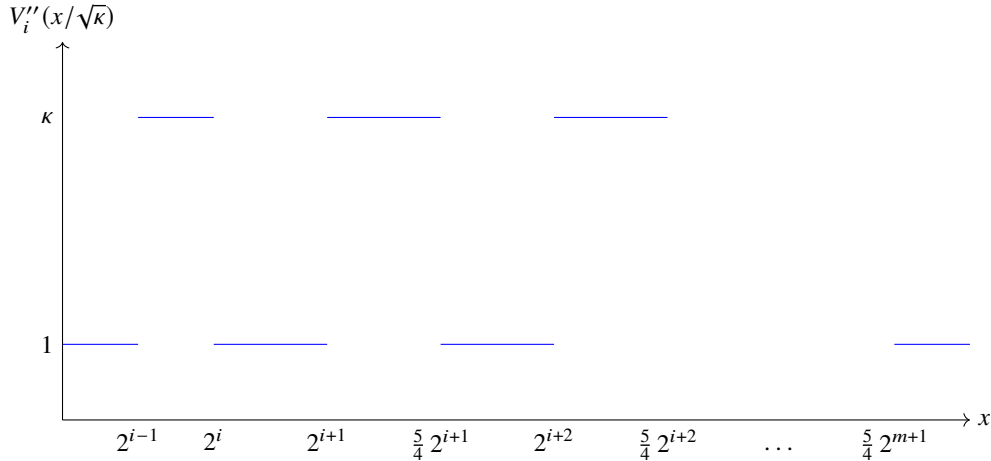


Figure 9.1 The horizontal axis is distorted for clarity.

There are two key properties of this construction. First, we will show in [Lemma 9.2.2](#) that each π_i places a substantial amount of mass on the interval $(\kappa^{-\frac{1}{2}} 2^{i-2}, \kappa^{-\frac{1}{2}} 2^{i-1}]$. This implies that if we can sample from π_i , it is likely that the sample will land in this interval, which is used to reduce the sampling task to the statistical testing task. Then, we will show in [Lemma 9.2.4](#) that V_i and V_{i+1} agree exactly outside of a small interval which is approximately located at $\kappa^{-\frac{1}{2}} 2^i$. This implies that for any given value $x \in \mathbb{R}^d$, there are only $O(1)$ possible values of $(V_i(x), V'_i(x))$ as i ranges in $[m]$, which in turn will be used to show that the oracle is not very informative (and hence prove a lower bound for the statistical testing task).

The intuition behind the following lemma is that at $\kappa^{-\frac{1}{2}} 2^{i-1}$, $V''_i = \kappa$ for the first time and so the density π_i drops off rapidly after this point.

Lemma 9.2.2. *For each $i \in [m]$,*

$$\pi_i((\kappa^{-\frac{1}{2}} 2^{i-2}, \kappa^{-\frac{1}{2}} 2^{i-1}]) \geq \frac{1}{32}.$$

Proof According to the definition of π_i , we have

$$\pi_i((\kappa^{-\frac{1}{2}} 2^{i-2}, \kappa^{-\frac{1}{2}} 2^{i-1}]) = \frac{\int_{\kappa^{-\frac{1}{2}} 2^{i-2}}^{\kappa^{-\frac{1}{2}} 2^{i-1}} \exp(-x^2/2) dx}{Z_{\pi_i}}, \quad Z_{\pi_i} := \int \exp(-V_i).$$

Recalling that m is chosen so that $\exp(-x^2/2) \geq 1/2$ whenever $|x| \leq \kappa^{-\frac{1}{2}} 2^{m-1}$,

$$\int_{\kappa^{-\frac{1}{2}} 2^{i-2}}^{\kappa^{-\frac{1}{2}} 2^{i-1}} \exp(-x^2/2) dx \geq \frac{1}{2} \kappa^{-\frac{1}{2}} 2^{i-2}.$$

For the normalizing constant, observe that

$$\int_0^\infty \exp(-V_i) = \int_0^{\kappa^{-\frac{1}{2}} 2^i} \exp(-V_i) + \int_{\kappa^{-\frac{1}{2}} 2^i}^\infty \exp(-V_i) \leq \kappa^{-\frac{1}{2}} 2^i + \int_{\kappa^{-\frac{1}{2}} 2^i}^\infty \exp(-V_i).$$

Since $V''_i = \kappa$ on $[\kappa^{-\frac{1}{2}} 2^{i-1}, \kappa^{-\frac{1}{2}} 2^i]$, it follows that $V'_i(\kappa^{-\frac{1}{2}} 2^i) \geq \kappa^{-\frac{1}{2}} 2^{i-1}$, and so

$$V_i(x) \geq \kappa^{-\frac{1}{2}} 2^{i-1} (x - \kappa^{-\frac{1}{2}} 2^i) + \frac{(x - \kappa^{-\frac{1}{2}} 2^i)^2}{2}, \quad x \geq \kappa^{-\frac{1}{2}} 2^i.$$

Therefore,

$$\int_{\kappa^{-\frac{1}{2}} 2^i}^\infty \exp(-V_i) \leq \int_{\kappa^{-\frac{1}{2}} 2^i}^\infty \exp\left(-\kappa^{-\frac{1}{2}} 2^{i-1} (x - \kappa^{-\frac{1}{2}} 2^i) - \frac{(x - \kappa^{-\frac{1}{2}} 2^i)^2}{2}\right) dx \leq \frac{1}{\kappa^{-\frac{1}{2}} 2^{i-1}} \leq \frac{1}{\sqrt{\kappa}},$$

where we applied a standard tail estimate for Gaussian densities ([Lemma 9.2.3](#)). Then,

$$\pi_i((\kappa^{-\frac{1}{2}} 2^{i-2}, \kappa^{-\frac{1}{2}} 2^{i-1}]) \geq \frac{2^{i-3}}{2(2^i + 1)} \geq \frac{1}{32},$$

which proves the result. $\square$

In the above proof, we used the following elementary lemma.

Lemma 9.2.3. Let $a, x_0 > 0$. Then,

$$\int_{x_0}^{\infty} \exp(-a(x - x_0) - \frac{1}{2}(x - x_0)^2) dx \leq \int_0^{\infty} \exp(-ax) dx = \frac{1}{a}.$$

The next lemma is the main reason why we used an oscillating construction for V_i'' .

Lemma 9.2.4. We have the equalities

$$V_i = V_{i+1}, \quad V_i' = V_{i+1}', \quad V_i'' = V_{i+1}'',$$

outside of the set $\{x \in \mathbb{R} : \kappa^{-\frac{1}{2}} 2^{i-1} \leq |x| \leq \frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}\}$.

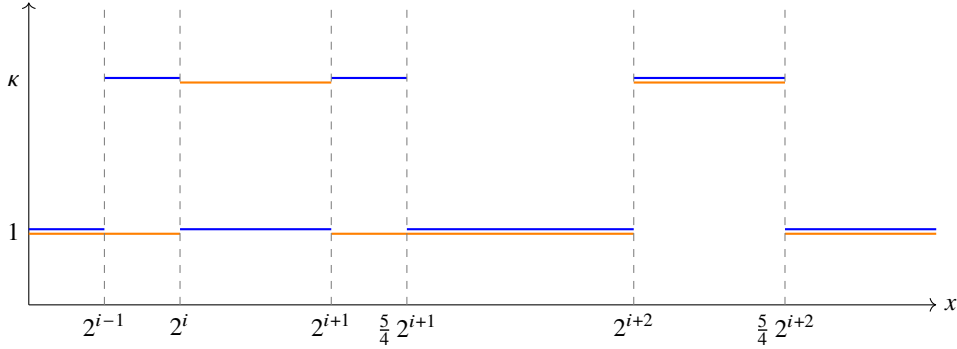


Figure 9.2 We plot $V_i''(x/\sqrt{\kappa})$ (in blue) and $V_{i+1}''(x/\sqrt{\kappa})$ (in orange). In this figure, we do not distort the horizontal axis lengths to make it easier to visually compare the relative lengths of intervals on which the second derivatives are constant.

Proof Refer to Figure 9.2 for a visual aid for the proof.

Clearly the potentials and derivatives match when $|x| \leq \kappa^{-\frac{1}{2}} 2^{i-1}$. Since the second derivatives match when $|x| \geq \frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}$, it suffices to show that

$$V_i'(\frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}) = V_{i+1}'(\frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}) \quad \text{and} \quad V_i(\frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}) = V_{i+1}(\frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}).$$

To that end, note that for $x \geq 0$,

$$V_{i+1}''(x) - V_i''(x) = \begin{cases} -(\kappa - 1), & \kappa^{-\frac{1}{2}} 2^{i-1} \leq x \leq \kappa^{-\frac{1}{2}} 2^i, \\ +(\kappa - 1), & \kappa^{-\frac{1}{2}} 2^i \leq x \leq \kappa^{-\frac{1}{2}} 2^{i+1}, \\ -(\kappa - 1), & \kappa^{-\frac{1}{2}} 2^{i+1} \leq x \leq \frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}, \\ 0, & \text{otherwise.} \end{cases}$$

A little algebra shows that the above expression integrates to zero, hence we deduce the equality

$V'_i(\frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}) = V'_{i+1}(\frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1})$. Also, by integrating this expression twice,

$$\begin{aligned}
 V_{i+1}(\frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}) - V_i(\frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}) &= \underbrace{-\frac{\kappa-1}{2} (\kappa^{-\frac{1}{2}} 2^{i-1})^2}_{\text{integral on } [\kappa^{-\frac{1}{2}} 2^{i-1}, \kappa^{-\frac{1}{2}} 2^i]} \\
 &\quad - \underbrace{(\kappa-1) \kappa^{-\frac{1}{2}} 2^{i-1} \kappa^{-\frac{1}{2}} 2^i + \frac{\kappa-1}{2} (\kappa^{-\frac{1}{2}} 2^i)^2}_{\text{integral on } [\kappa^{-\frac{1}{2}} 2^i, \kappa^{-\frac{1}{2}} 2^{i+1}]} \\
 &\quad + \underbrace{(\kappa-1) \kappa^{-\frac{1}{2}} 2^{i-1} \frac{1}{4} \kappa^{-\frac{1}{2}} 2^{i+1} - \frac{\kappa-1}{2} (\frac{1}{4} \kappa^{-\frac{1}{2}} 2^{i+1})^2}_{\text{integral on } [\kappa^{-\frac{1}{2}} 2^{i+1}, \frac{5}{4} \kappa^{-\frac{1}{2}} 2^{i+1}]} \\
 &= \frac{\kappa-1}{\kappa} \{-2^{2i-3} - 2^{2i-1} + 2^{2i-1} + 2^{2i-2} - 2^{2i-3}\} \\
 &= 0. \quad \square
 \end{aligned}$$

We need one final ingredient: **Fano's inequality**, which is the standard tool for establishing information-theoretic lower bounds.

Theorem 9.2.5 (Fano's inequality). *Let $\mathbf{i} \sim \text{uniform}([m])$. Then, for any estimator $\widehat{\mathbf{i}}$ of $\mathbf{i}$, where $\widehat{\mathbf{i}}$ is measurable with respect to some data Y ,*

$$\mathbb{P}\{\widehat{\mathbf{i}} \neq \mathbf{i}\} \geq 1 - \frac{\mathsf{l}(\mathbf{i}; Y) + \log 2}{\log m},$$

where l is the **mutual information** $\mathsf{l}(\mathbf{i}; Y) := \text{KL}(\text{law}(\mathbf{i}, Y) \parallel \text{law}(\mathbf{i}) \otimes \text{law}(Y))$.

Proof Let $H(\cdot)$ denote the **entropy** of a discrete random variable, i.e., if X has law p on a discrete alphabet $\mathcal{X}$, then $H(X) = \sum_{x \in \mathcal{X}} p(x) \log(1/p(x))$. We refer to [Cover and Thomas \(2006, §2\)](#) and [Exercise 9.1](#) for the basic properties of entropy (and related quantities).

Let $E := \mathbb{1}\{\widehat{\mathbf{i}} \neq \mathbf{i}\}$ denote the indicator of an error. Using the chain rule for entropy in two different ways,

$$\begin{aligned}
 H(\mathbf{i}, E \mid \widehat{\mathbf{i}}) &= H(\mathbf{i} \mid \widehat{\mathbf{i}}) + \underbrace{H(E \mid \mathbf{i}, \widehat{\mathbf{i}})}_{=0} \\
 &= H(E \mid \widehat{\mathbf{i}}) + H(\mathbf{i} \mid E, \widehat{\mathbf{i}}).
 \end{aligned}$$

Since conditioning reduces entropy, $H(E \mid \widehat{\mathbf{i}}) \leq H(E) \leq \log 2$. Also,

$$H(\mathbf{i} \mid E, \widehat{\mathbf{i}}) = \underbrace{\mathbb{P}\{\widehat{\mathbf{i}} = \mathbf{i}\} H(\mathbf{i} \mid \widehat{\mathbf{i}}, E = 0)}_{=0} + \mathbb{P}\{\widehat{\mathbf{i}} \neq \mathbf{i}\} H(\mathbf{i} \mid \widehat{\mathbf{i}}, E = 1) \leq \mathbb{P}\{\widehat{\mathbf{i}} \neq \mathbf{i}\} \log m.$$

Hence,

$$\mathbb{P}\{\widehat{\mathbf{i}} \neq \mathbf{i}\} \log m + \log 2 \geq H(\mathbf{i} \mid \widehat{\mathbf{i}}) = H(\mathbf{i}) - \mathsf{l}(\mathbf{i}; \widehat{\mathbf{i}}) \geq \log m - \mathsf{l}(\mathbf{i}; Y)$$

where the last inequality is the data-processing inequality ([Theorem 1.5.6](#)). Rearranging the inequality completes the proof of Fano's inequality. $\square$

Proof of Theorem 9.2.1, lower bound We follow the general outline described above.

1. Reduction to statistical testing. Let $\mathbf{i} \sim \text{uniform}([m])$ and suppose that for each $i \in [m]$, $\hat{\pi}_i$ is a distribution with $\|\hat{\pi}_i - \pi_i\|_{\text{TV}} \leq \frac{1}{64}$. Suppose that we have a sample $X \sim \hat{\pi}_{\mathbf{i}}$ (more precisely, this means that conditioned on $\mathbf{i} = i$, we have $X \sim \hat{\pi}_i$). In light of Lemma 9.2.2, a good candidate estimator $\hat{\mathbf{i}}$ for $\mathbf{i}$ is

$$\hat{\mathbf{i}} := i \in \mathbb{N} \quad \text{such that} \quad X \in (\kappa^{-\frac{1}{2}} 2^{i-2}, \kappa^{-\frac{1}{2}} 2^{i-1}] \quad \text{if such an } i \text{ exists.}$$

The probability that the estimator is correct is at least

$$\begin{aligned} \mathbb{P}\{\hat{\mathbf{i}} = \mathbf{i}\} &= \frac{1}{m} \sum_{i=1}^m \mathbb{P}\{\hat{\mathbf{i}} = i \mid \mathbf{i} = i\} = \frac{1}{m} \sum_{i=1}^m \mathbb{P}\{X \in (\kappa^{-\frac{1}{2}} 2^{i-2}, \kappa^{-\frac{1}{2}} 2^{i-1}] \mid \mathbf{i} = i\} \\ &= \frac{1}{m} \sum_{i=1}^m \hat{\pi}_i((\kappa^{-\frac{1}{2}} 2^{i-2}, \kappa^{-\frac{1}{2}} 2^{i-1}]) \geq \frac{1}{m} \sum_{i=1}^m \pi_i((\kappa^{-\frac{1}{2}} 2^{i-2}, \kappa^{-\frac{1}{2}} 2^{i-1}]) - \frac{1}{64} \geq \frac{1}{64}. \end{aligned} \quad (9.2.6)$$

Hence, a sampling can be used to solve the statistical testing problem.

2. A lower bound for the statistical testing problem. Next, we want to show for any algorithm which uses n queries to the oracle for π_i and outputs an estimator $\hat{\mathbf{i}}$ of $\mathbf{i}$, there is a lower bound for the probability of error $\mathbb{P}\{\hat{\mathbf{i}} \neq \mathbf{i}\}$.

First, suppose that the algorithm is *deterministic*, i.e., we assume that each query point x_j of the algorithm is a deterministic function of the previous query points and query values. Let $\mathcal{O}_i(x) := (V_i(x), V'_i(x), V''_i(x))$ denote the output of the oracle on input x when the target is π_i . Since the estimator $\hat{\mathbf{i}}$ is a function of $\{x_j, \mathcal{O}_i(x_j)\}_{j \in [n]}$, then Fano's inequality (Theorem 9.2.5) yields

$$\mathbb{P}\{\hat{\mathbf{i}} \neq \mathbf{i}\} \geq 1 - \frac{I(\mathbf{i}; \{x_j, \mathcal{O}_i(x_j)\}_{j \in [n]}) + \log 2}{\log m}.$$

By the chain rule for mutual information,

$$I(\mathbf{i}; \{x_j, \mathcal{O}_i(x_j)\}_{j \in [n]}) = \sum_{j=1}^n I(\mathbf{i}; x_j, \mathcal{O}_i(x_j) \mid \{x_{j'}, \mathcal{O}_i(x_{j'})\}_{j' \in [j-1]}). \quad (9.2.7)$$

By our assumption, conditioned on $\{x_{j'}, \mathcal{O}_i(x_{j'})\}_{j' \in [j-1]}$, the query point x_j is deterministic. Also, for a fixed point x_j , Lemma 9.2.4 implies that $\mathcal{O}_i(x_j)$ can only take on a constant number of possible values as i ranges over $[m]$ (the careful reader can check that the number of possible values for $\mathcal{O}_i(x_j)$ is at most 5). Together with (9.2.7),

$$I(\mathbf{i}; \{x_j, \mathcal{O}_i(x_j)\}_{j \in [n]}) \leq n \log 5.$$

Fano's inequality then yields

$$\mathbb{P}\{\hat{\mathbf{i}} \neq \mathbf{i}\} \geq 1 - \frac{n \log 5 + \log 2}{\log m}. \quad (9.2.8)$$

In general, for a possibly randomized algorithm, we can still deduce (9.2.8) by applying the previous argument conditioned on the random seed of the algorithm (which is independent of $\mathbf{i}$).

3. Finishing the argument. By combining together (9.2.6) and (9.2.8), and recalling that $m = \Theta(\log \kappa)$, we have shown that $n \gtrsim \log \log \kappa$. $\square$

The argument above makes rigorous the following intuition: since there are m distributions in our lower bound construction, there are $\log_2 m$ bits of information to learn. On the other hand, [Lemma 9.2.4](#) implies that each oracle query only reveals $O(1)$ bits of information. Hence, the number of queries required is at least $\Omega(\log m) = \Omega(\log \log \kappa)$.

Upper bound.

To show that the lower bound is tight, we develop an algorithm, based on rejection sampling, which achieves the lower complexity bound. As per our discussion in [Section 7.1](#), to implement rejection sampling we must specify the construction of an upper envelope $\tilde{\mu} \geq \tilde{\pi}$, where $\pi \in \Pi_\kappa$.

Without loss of generality, we assume that $V(0) = 0$ (if not, replace the output $V(x)$ of an oracle query with $V(x) - V(0)$). The upper bound algorithm only requires a zeroth-order oracle, and it is as follows.

- 1 Find the first index $i_- \in \{0, 1, \dots, \lceil \frac{1}{2} \log_2 \kappa \rceil\}$ such that $V(-2^{i_-}/\sqrt{\kappa}) \geq \frac{1}{2}$.
- 2 Find the first index $i_+ \in \{0, 1, \dots, \lceil \frac{1}{2} \log_2 \kappa \rceil\}$ such that $V(+2^{i_+}/\sqrt{\kappa}) \geq \frac{1}{2}$.
- 3 Set $x_- := -2^{i_-}/\sqrt{\kappa}$ and $x_+ := +2^{i_+}/\sqrt{\kappa}$; then, set

$$\tilde{\mu}(x) := \begin{cases} \exp\left(-\frac{x-x_-}{2x_-} - \frac{(x-x_-)^2}{2}\right), & x \leq x_-, \\ 1, & x_- \leq x \leq x_+, \\ \exp\left(-\frac{x-x_+}{2x_+} - \frac{(x-x_+)^2}{2}\right), & x \geq x_+. \end{cases}$$

To see why i_- and i_+ exist, from $V'' \geq 1$ and $V(0) = V'(0) = 0$ we have $V(x) \geq x^2/2$. Hence, if $|x| = 2^i/\sqrt{\kappa}$ where $i \geq \frac{1}{2} \log \kappa$, we have $V(x) \geq 1/2$.

Since V is decreasing (resp. increasing) on $\mathbb{R}_-$ (resp. $\mathbb{R}_+$), the first two steps can be implemented by running binary search over arrays of size $O(\log \kappa)$, which therefore only requires $O(\log \log \kappa)$ queries. We will prove that $\tilde{\mu}$ is a valid upper envelope for the unnormalized target $\tilde{\pi} := \exp(-V)$, and that $Z_\mu/Z_\pi \leq 1$. In turn, [Theorem 7.1.1](#) shows that once $\tilde{\mu}$ is constructed, an exact sample can be drawn from π using $O(1)$ additional queries in expectation.

Alternatively, if we require that the algorithm use a fixed (non-random) number of iterations, then note that in order to make the failure probability (the probability that rejection sampling fails to terminate within the allotted number of iterations) at most ε , it suffices to run rejection sampling for $O(\log(1/\varepsilon))$ steps. Combining this with the cost of constructing $\tilde{\mu}$, we conclude that we can output a sample whose law is ε -close to π in total variation distance using $O(\log \log \kappa + \log(1/\varepsilon))$ queries.

Proof of Theorem 9.2.1, upper bound First, we prove that $\tilde{\mu}$ is a valid upper envelope. Since $\tilde{\pi}$ is decreasing on $\mathbb{R}_+$ with $\tilde{\pi}(0) = \exp(-V(0)) = 1$, then $\tilde{\pi} \leq 1 \leq \tilde{\mu}$ on $[0, x_+]$. Next, since $V(x_+) \geq 1/2$ (by the definition of x_+), convexity of V yields

$$V'(x_+) \geq \frac{V(x_+) - V(0)}{x_+} \geq \frac{1}{2x_+}.$$

Thus, for $x \geq x_+$,

$$V(x) \geq V(x_+) + V'(x_+) (x - x_+) + \frac{1}{2} (x - x_+)^2 \geq \frac{1}{2x_+} (x - x_+) + \frac{1}{2} (x - x_+)^2,$$

which shows that $\tilde{\pi}(x) \leq \tilde{\mu}(x)$. By a symmetric argument on $\mathbb{R}_-$, we conclude that $\tilde{\pi} \leq \tilde{\mu}$.

By [Theorem 7.1.1](#), it suffices to bound Z_μ/Z_π . First, we claim that $\int_0^{x_+} \tilde{\mu} \gtrsim x_+$. When $i_+ = 0$, this holds

$$\int_0^{x_+} \tilde{\mu} = \int_0^{1/\sqrt{\kappa}} \exp(-V) \geq \int_0^{1/\sqrt{\kappa}} \exp\left(-\frac{\kappa x^2}{2}\right) dx \geq \frac{1}{3\sqrt{\kappa}} = \frac{x_+}{3}.$$

When $i_+ > 0$, then by the definition of i_+ we have $V(x_+/2) \leq 1/2$, so

$$\int_0^{x_+} \tilde{\mu} \geq \int_0^{x_+/2} \exp(-V) \geq \frac{x_+}{4}.$$

On the other hand, by [Lemma 9.2.3](#),

$$\int_{\mathbb{R}_+} \tilde{\mu} = \int_0^{x_+} \tilde{\mu} + \int_{x_+}^{\infty} \tilde{\mu} \leq x_+ + \int_{x_+}^{\infty} \exp\left(-\frac{1}{2x_+}(x - x_+) - \frac{1}{2}(x - x_+)^2\right) dx \leq 3x_+.$$

Hence, $\int_{\mathbb{R}_+} \tilde{\mu} \leq 3x_+ \leq 12 \int_{\mathbb{R}_+} \tilde{\pi}$, and similarly $\int_{\mathbb{R}_-} \tilde{\mu} \leq 12 \int_{\mathbb{R}_-} \tilde{\pi}$. Therefore, $Z_\mu/Z_\pi \leq 12$. $\square$

The obtained complexity $\Theta(\log \log \kappa)$ is surprisingly small; in particular, the upper bound uses a tailor-made algorithm based on rejection sampling, rather than any of the other existing algorithms (such as those based on Langevin dynamics). This is perhaps the best-case scenario for a lower bound: it helps us to determine if our existing algorithms are optimal, and if not, it gives guidance on how to design a better one.

9.3 Query Complexity in Constant Dimension

What about higher dimension? In particular, consider

$$\Pi_{\kappa,d} := \{\pi \in \mathcal{P}_{\text{ac}}(\mathbb{R}^d) \mid \pi \propto \exp(-V), I_d \leq \nabla^2 V \leq \kappa I_d, \nabla V(0) = 0\}.$$

Theorem 9.3.1 ([Chewi et al. \(2023b\)](#)). *There exists a universal constant $\varepsilon_0 > 0$ such that for any $d \geq 2$, $\mathcal{C}(\Pi_{\kappa,d}, \|\cdot\|_{\text{TV}}, \varepsilon_0) \asymp_d \log \kappa$.*

This result characterizes the *low-dimensional* complexity of log-concave sampling. The lower bound is based on a two-dimensional construction witnessing $\mathcal{C}(\Pi_{\kappa,2}, \|\cdot\|_{\text{TV}}, \varepsilon_0) \gtrsim \log \kappa$, which implies the lower bound in any dimension $d \geq 2$ because we can freely embed lower-dimensional distributions into higher-dimensional space (take the product with a Gaussian, for instance). On the other hand, the lower bound is tight in constant dimension because there is an algorithm witnessing $\mathcal{C}(\Pi_{\kappa,d}, \|\cdot\|_{\text{TV}}, \varepsilon_0) \lesssim_d \log \kappa$ for any d (here, the notation $\lesssim_d$ hides a large dependence on d).

Since the construction is rather technical, we will not present this result in detail and instead proceed on a more intuitive level.

Lower bound.

Let us revisit the intuition behind the one-dimensional construction. Due to [Lemma 9.2.2](#), π_i concentrates its mass on the interval $\kappa^{-1/2} [-2^{i-1}, +2^{i-1}]$. Therefore, consider instead the idealization $\tilde{\pi}_i = \text{uniform}(\kappa^{-1/2} [-2^{i-1}, +2^{i-1}])$ for $i \in [m]$. This family has the following desirable properties: (1) each $\tilde{\pi}_i$ is log-concave; (2) the distributions are well-separated, in the sense that $\|\tilde{\pi}_i - \tilde{\pi}_j\|_{\text{TV}} \gtrsim 1$ for $i, j \in [m]$, $i \neq j$ (this follows from [Lemma 9.2.2](#) and it is the reason why the intervals are chosen to double in size as i increases); (3) querying a distribution $\tilde{\pi}_i$ only yields a single bit of

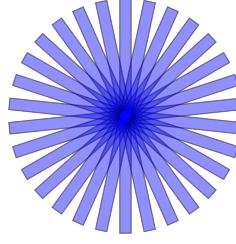


Figure 9.3 Uniform distributions over thin rectangles emanating from the origin.

information (whether or not the query lands in the support of $\tilde{\pi}_i$). Moreover, since a distribution in $\Pi_{\kappa,1}$ is “effectively supported” on an interval of length at least $\Omega(\kappa^{-1/2})$ and at most $O(1)$ (think about the upper bound algorithm from Section 9.2), we can take $m = \Theta(\log \kappa)$ and still have the family $\{\tilde{\pi}_i\}_{i \in [m]}$ “resemble” members of $\Pi_{\kappa,1}$. Since the number of queries needed to identify a member of the family is the logarithm of the size of the family, i.e., $\Theta(\log \log \kappa)$, this predicts the correct lower bound in dimension one. Of course, the fatal flaw is that the family $\{\tilde{\pi}_i\}_{i \in [m]}$ is *not* a subset of $\Pi_{\kappa,1}$, being neither strongly log-concave nor log-smooth, and most of the work in the construction is to find a true subset of $\Pi_{\kappa,1}$ that suitably mimics the idealized family $\{\tilde{\pi}_i\}_{i \in [m]}$.

Now we move to two dimensions. Here, our model for an element of $\Pi_{\kappa,2}$ is the uniform distribution over a rectangle of dimensions $\kappa^{-1/2} \times 1$. If we place these rectangles pointing away from the origin and spreading out like a fan, we can now fit $\Theta(\kappa^{1/2})$ of them in $\mathbb{R}^2$ while remaining well-separated (Figure 9.3). This family of uniform distributions still satisfies the three points listed above, and now we expect a query lower bound of $\Omega(\log \kappa)$ due to the larger size of the family.

The challenge is to now modify this family so that it lies in $\Pi_{\kappa,2}$, which is unfortunately quite involved. The first step is to note that strong log-concavity does not pose an issue: if we have a family of potentials which are convex and κ -smooth, then we can add a quadratic term $\kappa^{-O(1)} \|\cdot\|^2/2$ to each of the potentials, where $\kappa^{-O(1)}$ is chosen so that this term has a negligible effect in the region of interest. The condition number now becomes $\kappa^{O(1)}$, but this polynomial blow-up is acceptable since we are shooting for a $\Omega(\log \kappa)$ lower bound (it only affects the bound by a constant factor).

Let us restrict to rectangles in the positive quadrant and index them via bit strings $b \in \{0, 1\}^N$ corresponding to slope $[b] := 0.b_1 \dots b_N$ (represented in binary). Thus, we have 2^N distributions (eventually, $N \asymp \log \kappa$), where the b -th distribution will be concentrated on the sector $\mathcal{Z}_b := \{(x, y) \in \mathbb{R}^2 \mid x \geq 0, |y - [b]x| \leq 2^{-N}, \|(x, y)\| \leq 1\}$. In other words, $\mathcal{Z}_b$ will be the *zero set* of the potential V_b corresponding to the distribution $\pi_b \propto \exp(-V_b)$. In Figure 9.3, we essentially took π_b to be uniform over $\mathcal{Z}_b$ (with the immaterial difference that we now work with sectors rather than rectangles), but now we need to smooth out these distributions. We can do this by letting V_b grow slowly away from $\mathcal{Z}_b$, or by mollifying V_b ; eventually we will need to do both. As we will see, however, the former is convenient for controlling information leakage.

Indeed, as a first attempt, suppose we let $V_b = \text{dist}(\cdot, \mathcal{Z}_b)$. Then, V_b is convex because $\mathcal{Z}_b$ is convex, but it is non-smooth. We can fix the latter problem later via smoothing, but the bigger problem is that

this distribution reveals too much information via queries. Indeed, the gradient of V_b immediately reveals the direction toward $\mathcal{Z}_b$, and once we locate $\mathcal{Z}_b$ we can identify b and thus π_b . (Note that since our lower bound rests on finding a subfamily $\{\pi_b\}_{b \in \{0,1\}^N} \subseteq \Pi_{\kappa,2}$ that is already hard to sample, when proving our lower bound we have to contend with algorithms which “know” the family $\{\pi_b\}_{b \in \{0,1\}^N}$ in advance, hence identifying b suffices to output an exact sample from π_b .)

The main idea is to design the potential V_b so that querying at a distance of roughly 2^{-k} away from $\mathcal{Z}_b$ reveals only the first k bits of b . If this property holds, then heuristically the “probability” of learning k bits is only 2^{-k} , so in expectation we only learn $O(1)$ bits per query. Moreover, we can ensure this property as follows: let $[b]_k := 0.b_1 \dots b_k$ denote the slope corresponding to the first k bits of b , and define

$$\mathcal{Z}_{b,k} := \{(x, y) \in \mathbb{R}^2 \mid x \geq 0, |y - [b]_k x| \leq 2^{-k}, \|(x, y)\| \leq 1\}.$$

Then, we define

$$V_b(x, y) := \max_{k=1, \dots, N} 2^{-k} \text{dist}((x, y), \mathcal{Z}_{b,k}).$$

To see why this works, note that if we query at a large distance away from all of the $\mathcal{Z}_{b,k}$ ’s, then “the larger slope wins” in the maximum defining V_b , so $V_b = 2^{-1} \text{dist}(\cdot, \mathcal{Z}_{b,1})$: at this distance, the query to V_b only depends on the first bit of b . Similarly, convince yourself that a query at distance $\approx 2^{-k}$ away from $\mathcal{Z}_b$ only depends on $k + O(1)$ bits of b . Importantly, since V_b is a maximum of convex functions, it is convex.

Since V_b is still not smooth (its gradient is discontinuous), we still need to smooth further. For this, we can use the standard analysis trick of *mollification*, i.e., convolution with a smooth bump function which is supported on a ball of radius $\kappa^{-O(1)}$. After mollifying and adding a quadratic term as discussed above, all potentials are $\kappa^{-O(1)}$ -strongly convex and $\kappa^{O(1)}$ -smooth as desired. Unfortunately, mollification introduces a further complication: it “blurs” out the distribution, which causes information to “leak” between adjacent parts of the distribution. This is particularly harmful at the origin, where all of the zero sets intersect (Figure 9.3). As a result, it is conceivable that a single query near the origin, within the scale of the mollification, could reveal all of the bits of b at once.

Handling this last issue requires care, and this is addressed in Chewi et al. (2023b) with another multiscale construction. We refer to the paper for the full proof of the lower bound.

Upper bound.

Finally, to check that the lower bound is tight, we need to exhibit an algorithm with query complexity $O(\log \kappa)$ in constant dimension. Such a result is folklore in the literature on sampling from convex bodies, which shows how to “round” a convex body with a number of steps that is logarithmic in the condition number of the body (see, e.g., Kannan et al., 1997). We refer to Chewi et al. (2023b) for a self-contained analysis based on rejection sampling for our setting of interest, albeit with exponential dependence on d .

9.4 Query Complexity for Gaussians

The class of Gaussian distributions constitutes another setting in which we have a nearly complete understanding of the query complexity:

$$\Pi_{\kappa,d}^G := \{\text{normal}(0, \Sigma) : I_d \leq \Sigma^{-1} \leq \kappa I_d\}.$$

Theorem 9.4.1 (Chewi et al. (2023b)). *There exists a universal constant $\varepsilon_0 > 0$ such that the following statements hold.*

- 1 $\mathcal{E}(\Pi_{\kappa,d}^G, \|\cdot\|_{\text{TV}}, \varepsilon_0) \gtrsim d \wedge \sqrt{\kappa}$.
- 2 $\mathcal{E}(\Pi_{\kappa,d}^G, \|\cdot\|_{\text{TV}}, \varepsilon_0) \gtrsim \sqrt{\kappa} \log d$, provided $\kappa \leq d^{1/5-\delta}$ for some $\delta > 0$.
- 3 For any $\varepsilon > 0$, $\mathcal{E}(\Pi_{\kappa,d}^G, \sqrt{\text{KL}}, \varepsilon) \lesssim d \wedge \sqrt{\kappa} \log(d/\varepsilon^2)$.

These results state that, up to logarithmic factors, the query complexity for sampling from Gaussians is $d \wedge \sqrt{\kappa}$. The second statement above refines the first by adding the $\log d$ factor to the lower bound. We make the following remarks.

- Sampling from the class of Gaussians has nearly the same query complexity as optimization over the class of quadratic functions. In this setting, there is no gap between the two fields.
- In particular, the upper bound witnesses the accelerated rate of convergence (see the discussion in [Optimization Box 5.3.2](#)).
- The lower bound only exhibits a mild logarithmic dimension on the dependence, whereas all of the upper bounds for general log-concave sampling thus far incur polynomial dimension dependence. If we believe that this is intrinsic, then it suggests that Gaussians are not the hardest log-concave distributions to sample (unlike the case of optimization, where quadratics are essentially the hardest smooth convex functions to minimize).

The lower bounds borrow techniques from the literature on the complexity of linear algebraic tasks, measured in terms of the number of matrix-vector queries. Since the second lower bound is more involved, we only sketch the argument for the first lower bound.

Lower bound.

First, note that the potential for a Gaussian is quadratic: $V(x) = \frac{1}{2} \langle x, \Sigma^{-1} x \rangle$. Therefore, a first-order query yields $(V(x) - V(0), \nabla V(x)) = (\frac{1}{2} \langle x, \Sigma^{-1} x \rangle, \Sigma^{-1} x)$, where the second component gives strictly more information than the first. Hence, a query is of the form $x \mapsto \Sigma^{-1} x$.

We use a construction in which $\Lambda := \Sigma^{-1}$ is modelled as a Wishart distribution. Explicitly, we suppose that $\Lambda = XX^T$ where X is a $d \times d$ matrix with i.i.d. $\text{normal}(0, 1/d)$ entries. The lower bound is established via the following claims, which are left informal since we are only providing a sketch.

- **Claim 1.** With high probability, the eigenvalues of Λ lie in the range $[\Omega(d^{-2}), O(1)]$. Thus, the condition number of Λ is $\kappa \asymp d^2$.
- **Claim 2.** Suppose that there is a sampling algorithm which, with high probability, outputs a sample which is close in TV distance to $\pi = \text{normal}(0, \Lambda^{-1})$. Then, this sampler can be used to estimate $\text{tr}(\Lambda^{-1})$ up to a constant factor, with high probability.
- **Claim 3.** Any algorithm that estimates $\text{tr}(\Lambda^{-1})$ up to a constant factor, with high probability, must use $\Omega(d)$ matrix-vector queries for Λ .

These claims show that $\mathcal{E}(\Pi_{\kappa,d}^G, \|\cdot\|_{\text{TV}}, \varepsilon_0) \gtrsim d$ for $\kappa = cd^2$, where $c > 0$ is a universal constant. To handle other values of κ , we use a simple padding argument. Indeed, first suppose that $\kappa \geq cd^2$: in this case, $\mathcal{E}(\Pi_{\kappa,d}^G, \|\cdot\|_{\text{TV}}, \varepsilon_0) \geq \mathcal{E}(\Pi_{cd^2,d}^G, \|\cdot\|_{\text{TV}}, \varepsilon_0) \gtrsim d$. And in the other case $\kappa \leq cd^2$, since we can embed lower-dimensional examples into higher-dimensional space, $\mathcal{E}(\Pi_{\kappa,d}^G, \|\cdot\|_{\text{TV}}, \varepsilon_0) \geq \mathcal{E}(\Pi_{\kappa, \lfloor \sqrt{\kappa}/c \rfloor}^G, \|\cdot\|_{\text{TV}}, \varepsilon_0) \gtrsim \sqrt{\kappa}$. Hence, it suffices to establish the three claims above.

Claim 1. To put this fact into context, suppose that X is actually of size $n \times d$, again with i.i.d.

normal(0, 1/d) entries, and that $d, n \rightarrow \infty$ with aspect ratio $d/n \rightarrow \gamma$. Then, the limiting distribution of the spectrum of the Wishart matrix $XX^\top$ is governed by the Marchenko–Pastur law from random matrix theory. When $\gamma < 1$, i.e., $n < d$, the matrix $XX^\top$ is singular, so the condition number is infinite. When $\gamma > 1$, the matrix $XX^\top$ has full rank, and the limiting Marchenko–Pastur distribution has support bounded away from 0 and ∞ , suggesting that the condition number of $XX^\top$ should be $O(1)$.¹ On the other hand, here we are interested in the boundary case $d = n$ ($\gamma = 1$), where for any finite n the matrix $XX^\top$ is non-singular, but the condition number blows up as $d \rightarrow \infty$. It turns out that the blow-up is of order $\Theta(d^2)$ as claimed, but this is a non-trivial statement in random matrix theory; see [Edelman \(1989, Theorem 5.1\)](#) for results on the minimum eigenvalue.

Claim 2. Given a sample $Z \sim \text{normal}(0, \Lambda^{-1})$, note that $\mathbb{E}[\|Z\|^2] = \text{tr}(\Lambda^{-1})$. Thus, we can use samples from $\text{normal}(0, \Lambda^{-1})$, together with concentration of the norm, in order to estimate $\text{tr}(\Lambda^{-1})$. See [Exercise 9.3](#) for details.

Claim 3. This step is where we use the probabilistic structure of the Wishart matrix. The key fact is that if we condition a Wishart matrix on a sequence of n adaptive linear measurements, then up to an orthogonal transformation, then there is a block of the Wishart matrix that contains a “fresh” Wishart matrix of size $d - n$. The precise lemma is stated below.

Lemma 9.4.2 (Conditioning a Wishart matrix ([Braverman et al., 2020](#))). *Let Λ be a Wishart matrix of size $d \times d$. Then, conditional on any sequence of $n < d$ adaptive queries $x_1, \dots, x_n$ and the responses $\Lambda x_1, \dots, \Lambda x_n$, there exists an orthogonal matrix $U \in \mathbb{R}^{d \times d}$ and matrices $A \in \mathbb{R}^{n \times n}$, $B \in \mathbb{R}^{(d-n) \times n}$ depending only on the queries and responses such that*

$$U\Lambda U^\top = \begin{bmatrix} AA^\top & AB^\top \\ BA^\top & BB^\top + \tilde{\Lambda} \end{bmatrix},$$

where conditionally, $\tilde{\Lambda}$ has the Wishart distribution with size $(d - n) \times (d - n)$.

Now, the argument is concluded as follows. Suppose that we have made $n \leq d/2$ queries; then, the “unknown” portion of the Wishart matrix is another Wishart, with size $\Omega(d) \times \Omega(d)$. It turns out that the minimum eigenvalue of the Wishart has *heavy tails*, and this causes $\text{tr}(\Lambda^{-1})$ to have large fluctuations, even conditionally on the queries and responses. In the particular, for any estimator of $\text{tr}(\Lambda^{-1})$, there is a constant probability that the estimator will be off by more than a constant multiplicative factor. Hence, $\Omega(d)$ queries are needed to solve the inverse trace estimation problem. See [Exercise 9.4](#) for details of the argument.

Upper bound.

To show the matching upper bound, we use specific properties of the Gaussian; in particular, we can leverage the theory of polynomial approximation. See [Exercise 9.5](#).

Bibliographical Notes

The lower bounds in this chapter are from [Chewi et al. \(2021, 2023b\)](#). The algorithm for sampling from Gaussians in [Exercise 9.5](#) is closely related to the earlier work of [Nishimura and Suchard \(2023\)](#), based on the conjugate gradient method for solving linear systems. Below, we survey lower bounds in related settings and other approaches.

¹This can be made rigorous by studying the extreme eigenvalues.

The paper [Gopi et al. \(2022\)](#) obtained a zeroth-order query complexity result for distributions over the unit ball whose potentials V are of the form $V = \sum_{i=1}^{\infty} f_i + \frac{\alpha}{2} \|\cdot\|^2$, where each f_i is L -Lipschitz and the complexity is measured in terms of the number of queries to the individual f_i 's. In the regime $d \ll L^2/\alpha$, the query complexity is $\tilde{\Theta}(L^2/\alpha)$, where the upper bound is based on the proximal sampler (Section 8.1), and the lower bound is based on a reduction to optimization.

Lower bounds have also been obtained for stochastic gradient oracles; see the discussion in §10.1.1 and the corresponding Bibliographical Notes.

Other works have focused on obtaining lower bounds against specific families of algorithms. As already discussed in Section 7.3, the works [Chewi et al. \(2021\)](#); [Lee et al. \(2021a\)](#); [Wu et al. \(2021\)](#) proved lower bounds for MALA, culminating in a precise understanding of the runtime of MALA both from feasible and warm start initializations. On the other hand, the paper [Cao et al. \(2021\)](#) studied the complexity of discretization. In their setup, there is a stochastic process $(Z_t)_{t \geq 0}$, driven by some underlying Brownian motion $(B_t)_{t \geq 0}$; for example, the process $(Z_t)_{t \geq 0}$ could be the underdamped Langevin diffusion (Section 5.3). The algorithm is allowed to make queries to the potential V , as well as certain queries to the driving Brownian motion $(B_t)_{t \geq 0}$, and the goal of the algorithm is to output a point $\hat{Z}_T$ which is close to Z_T in mean squared error: $\mathbb{E}[\|Z_T - \hat{Z}_T\|^2] \leq \varepsilon^2$. Within this framework, they prove that the randomized midpoint discretization (Section 5.1) is optimal for simulating the underdamped Langevin dynamics (see [Theorem 5.3.12](#) for the upper bound). This framework, however, only provides lower bounds on the strong error; see [Srinivasan and Nagaraj \(2025\)](#).

Finally, another line of work considers the complexity of estimating the normalization constant (or partition function). In [Rademacher and Vempala \(2008\)](#), the authors studied the number of membership queries needed to estimate the volume of a convex body $K \subseteq \mathbb{R}^d$ such that $B(0, 1) \subseteq K \subseteq B(0, O(d^8))$ to within a small multiplicative constant; their lower bound for this problem is $\tilde{\Omega}(d^2)$. In comparison, the state-of-the-art upper bound for volume computation is $\tilde{O}(d^{3.5})$ (see [Cousins and Vempala, 2018](#); [Kook and Zhang, 2024](#); [Kook and Vempala, 2025c](#)).

In [Ge et al. \(2020\)](#), the authors considered the problem of estimating the normalizing constant $Z_\pi := \int \tilde{\pi}$ from queries to the unnormalized density $\tilde{\pi}$. Based on a multilevel Monte Carlo scheme, they showed that sampling algorithms can be turned into approximation algorithms for the normalizing constant, with the cost of an extra $O(d)$ dimension dependence in the reduction. By combining this with the upper bound from [Theorem 5.3.12](#), they showed that a $1 \pm \varepsilon$ multiplicative approximation to Z_π can be obtained using $\tilde{O}((d^{4/3}\kappa + d^{7/6}\kappa^{7/6})/\varepsilon^2)$ queries (in the strongly log-concave case). They then proved that $\Omega(d^{1-o(1)}/\varepsilon^{2-o(1)})$ queries are necessary. Unfortunately, due to the $O(d)$ loss in the reduction from estimating the normalizing constant to sampling, this does not imply a non-trivial lower bound for the task of sampling.

Exercises

Query Complexity in One Dimension

▷ **Exercise 9.1** (Basics of information theory)

Let X, Y be random variables taking values in finite alphabets $\mathcal{X}, \mathcal{Y}$ respectively. Let $p_{X,Y}$ denote the joint distribution, p_X denote the marginal distribution of X , etc. We make the following definitions.

- The entropy of X is $H(X) := \sum_{x \in \mathcal{X}} p_X(x) \log(1/p_X(x))$.
- The joint entropy of (X, Y) is $H(X, Y) := \sum_{x \in \mathcal{X}, y \in \mathcal{Y}} p_{X,Y}(x, y) \log(1/p_{X,Y}(x, y))$.
- The conditional entropy of X given Y is $H(X | Y) := \sum_{y \in \mathcal{Y}} p_Y(y) H(X | Y = y)$, where we define $H(X | Y = y) := \sum_{x \in \mathcal{X}} p_{X|Y}(x | y) \log(1/p_{X|Y}(x | y))$.

- The mutual information of X and Y is $I(X; Y) := \text{KL}(p_{X,Y} \parallel p_X \otimes p_Y)$.

Prove the following facts.

- 1 $0 \leq H(X) \leq \log |\mathcal{X}|$.
- 2 Chain rule for entropy: $H(X, Y) = H(X) + H(Y | X) = H(Y) + H(X | Y)$.
- 3 The mutual information satisfies $I(X; Y) = H(X) - H(X | Y) = H(Y) - H(Y | X)$.
- 4 Chain rule for mutual information: if $X_1, \dots, X_n$ are random variables taking values in $\mathcal{X}$, then $I(X_1, \dots, X_n; Y) = \sum_{i=1}^n I(X_i; Y | X_1, \dots, X_{i-1})$, where the conditional mutual information is given by $I(X_2; Y | X_1) = H(Y | X_1) - H(Y | X_1, X_2)$.

▷ **Exercise 9.2** (Sampling can be easier than distribution learning)

Here, we provide a simple example to show that sampling from a family of distributions can be much easier than learning from queries. Let p_0, p_1 be two densities on $[0, 1]$, defined as follows: $p_0(x) = 1 - 2x$ and $p_1(x) = 2x$. Consider the family of densities

$$\mathcal{P} := \left\{ \sum_{i=1}^N p_{a_i} \left(N \left(\cdot - \frac{i-1}{N} \right) \right) : a \in \{0, 1\}^N \right\}.$$

Suppose that we have access to a zeroth-order (or first-order) oracle for $\mathcal{P}$. Show that we can output an exact sample from a distribution in $\mathcal{P}$ in $O(1)$ queries. On the other hand, show that *learning* a distribution $p \in \mathcal{P}$, that is, outputting a density $\hat{p} : [0, 1] \rightarrow \mathbb{R}_+$ such that $\|\hat{p} - p\|_{\text{TV}} \leq c$ for some sufficiently small universal constant $c > 0$, requires $\Omega(N)$ queries.

This exercise demonstrates that when considering the difficulty of sampling, we must not be led astray by intuition concerning the difficulty of learning.

Query Complexity for Gaussians

▷ **Exercise 9.3** (Estimating the trace of the precision matrix)

Suppose that we can generate i.i.d. samples from a distribution $\hat{\pi}$ such that $\|\hat{\pi} - \pi\|_{\text{TV}} \leq \varepsilon_0$, where $\pi = \text{normal}(0, \Lambda^{-1})$. Show that for any $\delta > 0$, provided that $\varepsilon_0 \ll_\delta 1$ and using $O_\delta(1)$ samples from $\hat{\pi}$, one can output an estimator $\widehat{\text{tr}}$ such that $|\widehat{\text{tr}} - \text{tr}(\Lambda^{-1})| \leq \frac{1}{2} \text{tr}(\Lambda^{-1})$ with probability at least $1 - \delta$.

▷ **Exercise 9.4** (Lower bound for inverse trace estimation)

In this exercise, you may use the following tail bounds: for any $t > 0$ and $t' \in [0, 1]$,

$$\mathbb{P}\{\lambda_{\min}(\Lambda) \leq t'/d^2\} \asymp \sqrt{t'}, \quad \mathbb{P}\{\lambda_{\max}(\Lambda) \geq C(1+t)\} \leq 2 \exp(-dt),$$

and

$$\mathbb{P}\left\{\frac{1}{\lambda_j(\Lambda)} \geq \frac{d^2}{t^2}\right\} \leq (Ct)^{j^2},$$

where $C > 0$ is a universal constant. Suppose that $\widehat{\text{tr}}$ is an estimator based on n queries to Λ , such that with probability at least $1 - \delta$, $|\widehat{\text{tr}} - \text{tr}(\Lambda^{-1})| \leq \frac{1}{2} \text{tr}(\Lambda^{-1})$.

- 1 Show that there is a universal constant $C_1 > 0$ such that $\widehat{\text{tr}} \leq C_1 d^2$ with probability at least $1/2 - \delta$.
- 2 In the setting of [Lemma 9.4.2](#), show that $\lambda_{\min}(\Lambda) \leq \lambda_{\min}(\tilde{\Lambda})$.
- 3 Assume that $n \leq d/2$. Show that if δ is a sufficiently small universal constant, then $\widehat{\text{tr}} < \frac{1}{2} \text{tr}(\Lambda^{-1})$ with probability strictly bigger than δ . This contradiction shows that such an estimator cannot exist for $n \leq d/2$.

Hint: Lower bound this probability by

$$\mathbb{P}\{\widehat{\text{tr}} < \text{tr}(\Lambda^{-1})/2\} \geq \mathbb{P}\{\widehat{\text{tr}} \leq C_1 d^2 \text{ and } \lambda_{\max}(\tilde{\Lambda}^{-1}) > 2C_1 d^2\}.$$

Condition on the queries and responses, and use the tail bound for the minimum eigenvalue of a Wishart matrix.

- 4 Complete the proof of the first lower bound in [Theorem 9.4.1](#) by combining together this result with the result of the preceding exercise (conditioning on the event that the Wishart matrix has condition number $O(d^2)$).

▷ **Exercise 9.5** (An optimal algorithm for sampling from Gaussians)

Let $\pi = \text{normal}(0, \Sigma)$ where $I_d \leq \Sigma^{-1} \leq \kappa I_d$. You may use the following fact about polynomial approximation ([Chewi et al., 2023b](#)): for any $\delta \in (0, 1/2)$, there exists a polynomial p such that $\sup_{x \in [0, \kappa]} |p(x) - x^{-1/2}| \leq \delta/\sqrt{\kappa}$, where p has degree $O(\sqrt{\kappa} \log(\kappa/\delta))$. Show that this implies the upper bound in [Theorem 9.4.1](#).

Hint: Start with $X_0 \sim \text{normal}(0, I_d)$, and observe that by repeatedly querying the gradient oracle starting from X_0 , we can calculate $\Sigma^{-n} X_0$ for any $n \in \mathbb{N}$.

So far, we have only considered sampling within the black-box model, in which we only have access to oracle queries to the potential and its gradient. We will now consider several new sampling algorithms which go beyond the black-box model.

10.1 Stochastic Gradients

10.1.1 Bounded Variance

In this section, we study methods which use unbiased stochastic estimators of the gradient. Here, we begin with the simplest setting in which the gradient estimator is assumed to have uniformly bounded variance (or other suitable tail behavior). More precisely, assume that we have access to an oracle which, given $x \in \mathbb{R}^d$, returns a random vector $\hat{\nabla}V(x)$ such that

$$\mathbb{E} \hat{\nabla}V(x) = \nabla V(x), \quad \mathbb{E} [\|\hat{\nabla}V(x) - \nabla V(x)\|^2] \leq \sigma^2.$$

Implicitly, we are assuming that each query to the oracle yields a stochastic gradient that is *independent* of the query point and all previous stochastic gradients.

Naturally, we can consider substituting the stochastic gradient $\hat{\nabla}V$ into the Langevin diffusion, which leads to **stochastic gradient Langevin Monte Carlo (SG-LMC)**:

$$\hat{X}_{(n+1)h} := \hat{X}_{nh} - h \hat{\nabla}V(\hat{X}_{nh}) + \sqrt{2} (B_{(n+1)h} - B_{nh}). \quad (\text{SG-LMC})$$

A convenient approach to analyzing this iteration is to apply the local error framework in [Lemma 5.1.2](#). Indeed, if we let $\hat{X}_h, X_h$ denote the iterates of **SG-LMC** and the Langevin diffusion respectively, both started at x , then the unbiasedness of the stochastic gradient yields

$$\mathbb{E} \hat{X}_h = \mathbb{E}[x - h \hat{\nabla}V(x) + \sqrt{2} B_h] = x - h \nabla V(x),$$

and therefore

$$\|\mathbb{E} \hat{X}_h - \mathbb{E} X_h\| = \mathbb{E} \left\| \int_0^h \{\nabla V(X_t) - \nabla V(x)\} dt \right\|,$$

and hence the weak error of **SG-LMC** is the same as for **LMC**. On the other hand, by using the variance bound, the strong error picks up an additional term scaling as σh . Substituting this into [Lemma 5.1.2](#) yields an complexity of

$$N = \tilde{O}\left(\frac{\kappa^2 d}{\varepsilon^2} + \frac{\sigma^2}{\alpha \varepsilon^2}\right) \quad \text{iterations} \quad (10.1.1)$$

to achieve $\sqrt{\alpha} W_2(\hat{\mu}_{Nh}, \pi) \leq \varepsilon$, in our usual strongly convex and smooth setting; see [Exercise 5.1](#).

The first term in the bound above can be improved using the techniques in this book. For instance, by considering a stochastic gradient version of **RM-ULMC**, a slight modification of [Exercise 5.12](#) would yield a rate of $\tilde{O}(\kappa d^{1/3}/\varepsilon^{2/3} + \sigma^2/\alpha \varepsilon^2)$.

On the other hand, readers who have studied convex optimization may recognize the $\sigma^2/\alpha \varepsilon^2$ term: it is the rate achieved by stochastic gradient descent on a strongly convex loss¹, and in that setting, it is even optimal in an information-theoretic sense ([Raginsky and Rakhlin, 2011](#); [Agarwal et al., 2012](#)). However, it turns out to not be optimal for the problem of sampling.

We start by giving a simple lower bound from [Chen et al. \(2026a\)](#).

Theorem 10.1.2. *Let $\psi : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ be an increasing function with $\psi(0) = 0$, and let $\varepsilon > 0$. Suppose that there is an algorithm which, given access to a stochastic gradient oracle for $\pi \in \{\text{normal}(0, 1), \text{normal}(\varepsilon, 1)\}$, outputs a sample with TV distance at most $\varepsilon/10$ away from π . The stochastic gradient oracle is assumed to be unbiased, $\mathbb{E} \hat{V}'(x) = V'(x)$, and satisfy $\mathbb{E} \psi(|\hat{V}'(x) - V'(x)|) \leq 1$ for all $x \in \mathbb{R}$. Then, the algorithm must make at least N queries, where*

$$N \geq \frac{1}{10} \sup\{u \geq \varepsilon \mid \varepsilon \psi(u) \leq (1 - \psi(\varepsilon)) u\}.$$

Since the proof is relatively simple, it is given as [Exercise 10.2](#). Under the bounded variance assumption, we can take $\psi(u) = u^2/\sigma^2$. This shows that for $\varepsilon \ll \sigma$, $\Omega(\sigma^2/\varepsilon)$ queries are required to sample from the family of two Gaussian distributions. Moreover, in the general α -strongly convex case, we can consider the rescaled potentials $x \mapsto V(x/\sqrt{\alpha})$, which has the effect of replacing σ^2 by σ^2/α ; the resulting lower bound is therefore $\Omega(\sigma^2/\alpha \varepsilon)$. Notably, the dependence on $1/\varepsilon$ is substantially smaller than in the upper bounds above.²

In the same paper, [Chen et al. \(2026a\)](#) gave an algorithm with query complexity

$$N = \tilde{O}\left(\kappa d^{1/2} + \frac{\kappa d^{1/2} \sigma^2}{\alpha \varepsilon}\right)$$

in TV distance. Although the $1/\varepsilon$ dependence is optimal, the prefactor is not. Even more recently, [Paulin and Whalley \(2026\)](#) gave a refined analysis of a discretization of the underdamped Langevin diffusion which achieves a query complexity of

$$N = \tilde{O}\left(\frac{\kappa^{3/2} d^{1/2}}{\varepsilon} + \frac{\kappa^{1/2} \sigma^2}{\alpha \varepsilon}\right)$$

in W_1 distance, which holds under the extra assumption $\int \|\nabla \hat{V}(x)\|_{\text{op}}^2 \pi(\text{d}x) \leq \beta^2$. Although the

¹In optimization, the rate would typically be reported as $\sigma^2/\alpha \varepsilon$ since the goal is taken to be $V - \inf V \leq \varepsilon$, whereas in our convention we set the final error to be ε^2 .

²Technically, the bounds discussed in this section are somewhat incomparable as they hold in different metrics, but for simplicity we ignore this distinction. For example, the upper bound for **SG-LMC** could be improved to hold in KL divergence using the framework of [Altschuler and Chewi \(2026b\)](#).

story here is not yet settled, it seems likely that the query complexity using stochastic gradients in the bounded variance setting is the same as for the exact gradient case, with an extra $\sigma^2/\alpha\varepsilon$ term.

On the other hand, if the tails of the stochastic gradients decay faster, then the lower bound in [Theorem 10.1.2](#) also degrades. Suppose, for instance, that the stochastic gradients have sub-Gaussian tails: we take $\psi(u) = \exp(u^2/\sigma_{\psi_2}^2) - 1$, where $\sigma_{\psi_2}^2$ is the sub-Gaussian variance proxy. Then, the $1/\varepsilon$ dependence implied by [Theorem 10.1.2](#) is only $\Omega(\sqrt{\log(1/\varepsilon)})$. In this setting, [Chen et al. \(2026a\)](#) showed that high-accuracy guarantees for sampling are indeed achievable.

Theorem 10.1.3 ([Chen et al. \(2026a\)](#)). *There is an algorithm which, for any target distribution $\pi \propto \exp(-V)$ with $0 < \alpha I_d \leq \nabla^2 V \leq \beta I_d$ and $\nabla V(0) = 0$, when given access to a stochastic gradient oracle for V satisfying $\mathbb{E} \exp(\|\hat{\nabla} V(x) - \nabla V(x)\|^2 / \sigma_{\psi_2}^2) \leq 2$ for all $x \in \mathbb{R}^d$, outputs a sample from a distribution $\hat{\mu}$ with $\|\hat{\mu} - \pi\|_{\text{TV}} \leq \varepsilon$ using*

$$\tilde{O}\left(\left(\kappa d^{1/2} + \frac{\sigma_{\psi_2}^2}{\alpha}\right) \text{polylog} \frac{1}{\varepsilon}\right) \quad \text{queries}.$$

This result is perhaps surprising since stochasticity of the gradients precludes high-accuracy convergence for optimization, regardless of the tail behavior. However, in sampling, the stochasticity of the gradient can be “folded into” the inherent randomness of the algorithm.

10.1.2 Finite Sum

Often, stochastic gradients arise as mini-batches. Let us assume that the potential V is of finite-sum form, $V := m^{-1} \sum_{i=1}^m f_i$. For example, each f_i could correspond to the loss associated with a single data point, so that m is the size of the full data set. We assume that we can query the value and gradient of each f_i , and we measure the number of individual function queries. Thus, in this model, computing the full gradient ∇V costs $O(m)$ queries.

We assume that V is α -strongly convex and that each f_i is β -smooth; in particular, V is β -smooth. As always, we denote $\kappa := \beta/\alpha$.

The simplest approach is to draw a mini-batch of size B ([Welling and Teh, 2011](#)):

$$\hat{X}_{(n+1)h} := \hat{X}_{nh} - \frac{h}{B} \sum_{j=1}^B \nabla f_{i_n^j}(\hat{X}_{nh}) + \sqrt{2} (B_{(n+1)h} - B_{nh}), \quad i_n^1, \dots, i_n^B \stackrel{\text{i.i.d.}}{\sim} \text{uniform}([m]).$$

It is an instance of [SG-LMC](#) with a particular stochastic gradient estimator; however its variance cannot be controlled purely in terms of α , β , d , and m , even at stationarity, unless $B = m$ ([Exercise 10.3](#)). A simple remedy is to use a control variate ([Baker et al., 2019](#)): we add a centered term to the gradient estimate, thus retaining unbiasedness, but reducing the variance. Letting $B = 1$ for simplicity, consider

$$\hat{X}_{(n+1)h} := \hat{X}_{nh} - h \{\nabla f_{i_n}(\hat{X}_{nh}) - \nabla f_{i_n}(x_\star)\} + \sqrt{2} (B_{(n+1)h} - B_{nh}), \quad i_n \sim \text{uniform}([m]).$$

Here, $x_\star$ is the minimizer, and we precompute $\nabla f_i(x_\star)$ for each $i \in [m]$. The added term is indeed centered: $\mathbb{E} \nabla f_{i_n}(x_\star) = \nabla V(x_\star) = 0$. However, at stationarity $X \sim \pi$, the variance is now bounded:

$$\frac{1}{m} \sum_{i=1}^m \mathbb{E}[\|\nabla f_i(X) - \nabla f_i(x_\star)\|^2] \leq \beta^2 \mathbb{E}[\|X - x_\star\|^2] = \frac{\beta^2 d}{\alpha},$$

by Lemma 4.0.1. If we substitute this into (10.1.1), we expect the complexity bound

$$N = \tilde{O}\left(m + \frac{\kappa^2 d}{\varepsilon^2}\right) \quad \text{individual function queries.} \quad (10.1.4)$$

This is established rigorously via Lemma 5.1.2 (Exercise 10.4). Note that this is substantially better than LMC, which by Theorem 4.1.2 requires $\tilde{O}(\kappa^2 dm/\varepsilon^2)$ queries.³

Another further source of improvement comes from the family of techniques known as *variance reduction*. When we apply this to the Langevin diffusion, it turns out to not significantly improve upon (10.1.4). Namely, the bound (10.1.4) really consists of two sources of error: $\kappa^2 d/\varepsilon^2$ from discretization, and $\kappa^2 d/\varepsilon^2$ from the stochastic gradient. Note that these two terms balance; although variance reduction can reduce the latter, the overall bound is bottlenecked by the former. Thus, we would have to turn toward a more complicated algorithm as in Chapter 5 to see the benefit. Instead of doing so, we present the main idea of variance reduction in a simpler setting in Section 10.2.1.

We mention it here, however, because it leads to the current state-of-the-art results. When applied to RM-ULMC, Hu et al. (2021) obtained an algorithm using

$$N = \tilde{O}\left(m + \kappa^2 + \frac{\kappa^{4/3} d^{1/3} m^{2/3}}{\varepsilon^{2/3}}\right) \quad \text{individual function queries.}$$

Moreover, in the high-accuracy regime, the combined results of Lee et al. (2021c) and Chen et al. (2026a) yield an algorithm using

$$N = \tilde{O}\left((m + \kappa d^{1/2} m^{1/2}) \text{polylog}(1/\varepsilon)\right) \quad \text{individual function queries.}$$

These results lie beyond the scope of the book.

10.2 Coordinate Methods

Next, we consider methods which update individual coordinates or blocks of coordinates. In addition to yielding samplers with smaller per-iteration cost, such methods can sometimes take advantage of better directional smoothness along coordinate axes.

10.2.1 Coordinate Langevin

As a starting point, suppose that at each iteration n , we pick a random coordinate $i_n \in [d]$ and compute the partial derivative $\partial_{i_n} V$. We allow ourselves to choose the coordinates with different probabilities: let p_i denote the probability of choosing coordinate i . In order to keep the gradient estimator unbiased, we should then weight the partial derivative by a factor $p_{i_n}^{-1}$. This leads to the iteration

$$\hat{X}_{(n+1)h} = \hat{X}_{nh} - \frac{h}{p_{i_n}} \partial_{i_n} V(\hat{X}_{nh}) e_{i_n} + \sqrt{2} (B_{(n+1)h} - B_{nh}), \quad i_n \sim p. \quad (\text{CG-LMC})$$

Here, the acronym stands for **coordinate gradient Langevin Monte Carlo**. This is in fact a special case of SG-LMC with a particular gradient estimator.

³We know that the analysis of LMC can be improved to $\tilde{O}(\kappa dm/\varepsilon^2)$ by Theorem 4.3.6. However, for the sake of isolating the effect of mini-batches, we compare the bounds obtained purely from the local error coupling analysis.

Let us bound the variance of the stochastic gradient: for $i \sim p$,

$$\mathbb{E} \left[\left\| \frac{1}{p_i} \partial_i V(X) e_i - \nabla V(X) \right\|^2 \right] \leq \sum_{i=1}^d \frac{\mathbb{E}[|\partial_i V(X)|^2]}{p_i}.$$

At stationarity, for $X \sim \pi$, integration by parts shows that $\mathbb{E}[|\partial_i V(X)|^2] = \mathbb{E} \partial_i^2 V(X)$. Hence, if we assume the directional smoothness bounds

$$\partial_i^2 V(x) \leq \beta_i, \quad \text{for all } i \in [n] \text{ and } x \in \mathbb{R}^d,$$

the variance of the stochastic gradient is bounded (at stationarity) by $\sum_{i=1}^d \beta_i / p_i$. This is optimized by setting $p_i \propto \sqrt{\beta_i}$, with the optimal value being $(\sum_{i=1}^d \sqrt{\beta_i})^2$. Heuristically, the final rate is given by plugging in this variance bound into (10.1.1). By carrying out this argument rigorously (Exercise 10.5), we obtain the iteration complexity

$$N = \tilde{O} \left(\frac{\kappa^2 d + (\sum_{i=1}^d \sqrt{\kappa_i})^2}{\varepsilon^2} \right) \quad \text{partial derivative evaluations,} \quad (10.2.1)$$

for sufficiently small $\varepsilon > 0$, where $\kappa_i := \beta_i / \alpha$, and $\kappa := \beta / \alpha$ with β denoting the full smoothness (the largest operator norm of $\nabla^2 V$).

To understand this bound, suppose that the evaluation of the full gradient ∇V is d times as costly as the evaluation of a partial derivative $\partial_i V$. This is not the case for all applications, but it is a natural setting in which it is sensible to consider coordinate methods. Under this model, the cost of the **LMC** baseline is $\tilde{O}(\kappa^2 d^2 / \varepsilon^2)$.⁴ Now, we know that $\beta_i \leq \beta$ always, which means that the second term in the bound (10.2.1) is at most $\tilde{O}(\kappa^2 d^2 / \varepsilon^2)$: under this cost model, and based on our upper bounds, **CG-LMC** is never worse than **LMC**. On the other hand, it is possible to have $\beta_i \ll \beta$. The worst case, assuming that $\nabla^2 V \geq 0$, is when $\nabla^2 V$ looks like a multiple of the all-ones matrix, for which $\beta_i \asymp \beta / d$. In such a case, the bound 10.2.1 reads $\tilde{O}(\kappa^2 d / \varepsilon^2)$, substantially better than **LMC**.

The **CG-LMC** algorithm is not a pure coordinate method because it does not update a single coordinate at each step: it adds Brownian motion to every coordinate. This is also the reason why the bound (10.2.1) involves the full smoothness. An alternative approach is to use **random coordinate Langevin Monte Carlo**, as studied in Ding et al. (2021a):

$$\hat{X}_{(n+1)h} := \hat{X}_{nh} - \frac{h}{p_{i_n}} \partial_{i_n} V(\hat{X}_{nh}) e_{i_n} + \sqrt{\frac{2}{p_{i_n}}} \langle B_{(n+1)h} - B_{nh}, e_{i_n} \rangle e_{i_n}, \quad i_n \sim p. \quad (\text{RC-LMC})$$

To analyze this method, what is the ideal process to which we should compare **RC-LMC**? For each $i \in [d]$, let P_i denote the Markov kernel corresponding to running the following SDE for time h :

$$dX_t = -\frac{1}{p_i} \partial_i V(X_t) e_i dt + \sqrt{\frac{2}{p_i}} e_i e_i^\top dB_t. \quad (10.2.2)$$

Also, let $P := \sum_{i=1}^d p_i P_i$. We can check that each P_i leaves π invariant: for any test function ϕ , along (10.2.2), $\partial_t \mathbb{E} \phi(X_t) = p_i^{-1} \mathbb{E}[\partial_i^2 \phi(X_t) - \partial_i V(X_t) \partial_i \phi(X_t)]$, and this vanishes for $X_t \sim \pi$ by integration by parts. As noted in Ding et al. (2021a), however, there is an additional difficulty here, as one cannot simply use strong convexity of V to prove that the SDE (10.2.2) contracts.

There is a trick to make the analysis fit within the local error framework. Whereas Lemma 5.1.2 is typically applied with P being the kernel for the “ideal” process, and $\hat{P}$ the kernel of the discretization,

⁴As in Section 10.1.2, Theorem 4.3.6 improves this to $\tilde{O}(\kappa d^2 / \varepsilon^2)$, but we ignore this for the sake of comparison.

we can instead swap the roles of $\hat{P}$ and P . Thus, the W_2 contractivity assumption is now imposed upon $\hat{P}$, which makes it easy to verify. A careful computation of the local errors yields the iteration complexity bound (Exercise 10.6):

$$N = \tilde{O}\left(\frac{(\sum_{i=1}^d \kappa_i)^2}{\varepsilon^2}\right) \quad \text{partial derivative evaluations,} \quad (10.2.3)$$

for sufficiently small $\varepsilon > 0$. This bound turns out to be incomparable with (10.2.1).

The results can be further improved by incorporating the technique of variance reduction. Consider the following method, which combines the stochastic variance reduced gradient (SVRG) method of Johnson and Zhang (2013) with CG-LMC. Following Ding and Li (2021), we call it **SVRG-LMC**.

$$\begin{aligned} \hat{X}_{(n+1)h} &:= \hat{X}_{nh} - h \hat{\nabla}_n V + \sqrt{2} (B_{(n+1)h} - B_{nh}), \\ \hat{\nabla}_n V &:= \nabla V(\hat{X}_{\lfloor n/K \rfloor Kh}) + d \{ \partial_{i_n} V(\hat{X}_{nh}) - \partial_{i_n} V(\hat{X}_{\lfloor n/K \rfloor Kh}) \} e_{i_n}. \end{aligned} \quad (\text{SVRG-LMC})$$

For simplicity, we consider uniform weighting: $i_n \sim \text{uniform}([d])$. Here, K is the “epoch length”; note that we compute the full gradient ∇V once per epoch, and we compute one new partial derivative per iteration. The idea is that the full gradients help to significantly reduce the variance, but remain inexpensive as long as we do not compute them too frequently.

Let us denote by $\bar{X} := \hat{X}_{\lfloor n/K \rfloor Kh}$ the so-called anchor point. The variance is bounded as follows:

$$\begin{aligned} \mathbb{E}[\|\hat{\nabla}_n V - \nabla V(\hat{X}_{nh})\|^2] &\leq d \sum_{i=1}^d \mathbb{E}[|\partial_i V(\hat{X}_{nh}) - \partial_i V(\bar{X})|^2] = d \mathbb{E}[\|\nabla V(\hat{X}_{nh}) - \nabla V(\bar{X})\|^2] \\ &\leq \beta^2 d \mathbb{E}[\|\hat{X}_{nh} - \bar{X}\|^2]. \end{aligned}$$

The variance bound now depends on $\mathbb{E}[\|\hat{X}_{nh} - \bar{X}\|^2]$, which is controlled by the epoch length. In fact, by Lemma 4.1.8, we expect the Brownian motion to yield the dominant contribution, and heuristically $\mathbb{E}[\|\hat{X}_{nh} - \bar{X}\|^2] \lesssim Kdh$. If we now choose the epoch length $K \asymp d$, the error analysis leads to the step size choice $h = \tilde{\Theta}(\min\{\varepsilon^2/\beta\kappa d, \varepsilon/\beta d^{3/2}\})$. The total number of iterations is $N_{\text{iter}} = \tilde{\Omega}(1/\alpha h)$, and the total number of partial derivative evaluations is $O(N_{\text{iter}}(1 + m/K)) = O(N_{\text{iter}})$. The total cost is

$$N = \tilde{O}\left(\frac{\kappa^2 d}{\varepsilon^2} + \frac{\kappa d^{3/2}}{\varepsilon}\right) \quad \text{partial derivative evaluations.} \quad (10.2.4)$$

See Exercise 10.7.

Note that in our partial derivative computational model, the $d^{3/2}$ dimension dependence is equivalent to the results for **RM-LMC** and **ULMC**. This may seem surprising, as we are still using a variant of the basic **LMC** algorithm. The explanation is that the discretization error term remains the same, but we have reduced the per-iteration cost. So, we will see a net benefit as long as we can keep the stochastic gradient error controlled, and this is achieved via variance reduction.

Pushing this further, Ding and Li (2021) combined variance reduction with **ULMC**, yielding algorithms which use

$$N = \tilde{O}\left(\frac{\kappa^{3/2} d^{1/2}}{\varepsilon} + \frac{\kappa d^{4/3}}{\varepsilon^{2/3}}\right) \quad \text{partial derivative evaluations.} \quad (10.2.5)$$

Plausibly, the first term in this bound could be further improved by replacing **ULMC** with **RM-ULMC**. Finally, we mention that in the high-accuracy regime, Lu and Wang (2022a) showed that the zigzag sampler, from a warm start, succeeds using

$$N = \tilde{O}(\kappa^2 d^{3/2} \text{polylog}(1/\varepsilon)) \quad \text{partial derivative evaluations.}$$

10.2.2 Gibbs Sampling

In this section, we analyze the classical Gibbs sampling algorithm, following the remarkable work of [Ascolani et al. \(2024\)](#). We assume that the target distribution is of the form

$$\pi(x) \propto \exp\left(-V_0(x) - \sum_{i=1}^M V_i(x^i)\right), \quad (10.2.6)$$

where $x := x^{1:M} := (x^1, \dots, x^M)$. We allow the x^i to be blocks of coordinates: $x^i \in \mathbb{R}^{d_i}$, $\sum_{i=1}^M d_i = d$. The **random scan Gibbs sampling (RS-GS)** performs the update

$$X_{n+1}^{i_n} \sim \exp(-V_0(X_n^{-i_n}, \cdot) - V_{i_n}(\cdot)), \quad i_n \sim \text{uniform}([M]), \quad (\text{RS-GS})$$

where the other blocks are left unchanged. Here, $X_n^{-i_n} := (X_n^1, \dots, X_n^{i_n-1}, X_n^{i_n+1}, \dots, X_n^M)$, and we abuse notation by writing $V_0(X_n^{-i_n}, \cdot)$ where it is understood that the arguments are passed into V_0 in the appropriate order.

Gibbs sampling assumes that we can sample from the conditional distributions $\pi^{i|-i}$ exactly. A typical application is when each block is of size 1 ($d_i = 1$), in which case the conditional distributions are univariate and presumably easy to sample. Regarding the form of (10.2.6), we could have simply taken $V_i = 0$ for all $i \geq 1$, which is our usual model. However, it turns out that the addition of the block-separable term does not hinder the analysis, so we include it for greater generality.

We can express **RS-GS** in terms of Markov kernels as follows: for each $i \in [M]$ and $x \in \mathbb{R}^d$, let $P_i(x, \cdot) := \pi^{i|-i}(\cdot | x^{-i})$. Then, **RS-GS** iterates the Markov kernel $P := M^{-1} \sum_{i=1}^M P_i$.

Theorem 10.2.7 (Random scan Gibbs sampling, [Ascolani et al. \(2024\)](#)). *We assume:*

- $V_0(x^{-i}, \cdot)$ is β_i -smooth for all $i \in [M]$ and $x^{-i} \in \mathbb{R}^{d-d_i}$.
- V_0 is κ^{-1} -strongly convex relative to the norm $\|\cdot\|_\beta$, where $\|x\|_\beta^2 := \sum_{i=1}^M \beta_i \|x^i\|^2$.
- V_i is convex for each $i \in [M]$.

Then, for any $\mu \in \mathcal{P}(\mathbb{R}^d)$,

$$\text{KL}(\mu P \parallel \pi) \leq \left(1 - \frac{1}{\kappa M}\right) \text{KL}(\mu \parallel \pi). \quad (10.2.8)$$

First note that when $\beta_i = \beta$ for all $i \in [M]$, then κ is simply the usual condition number of V_0 , but the definition above allows for anisotropic smoothness as in the discussion of Section 10.2.1. Also, in the standard case $d_i = 1$ for all i , then $M = d$, so the iteration complexity implied by [Theorem 10.2.7](#) is $\tilde{O}(\kappa d \log(1/\varepsilon))$. In some situations, since each iteration of **RS-GS** involves sampling from a univariate distribution, the per-iteration cost could be considered to be $O(1)$, whereas an evaluation of the gradient ∇V_0 could be considered to have cost $O(d)$. Under this cost model, the guarantee for **RS-GS** would be far better than the guarantees for other methods such as **LMC**. Although reasonable, whether or not this cost model is appropriate depends on the specific application at hand.

The rest of this section is devoted to the proof of [Theorem 10.2.7](#).

Proof of Theorem 10.2.7 The first step is to note that by convexity and the chain rule,

$$\text{KL}(\mu P \parallel \pi) = \text{KL}\left(\frac{1}{M} \sum_{i=1}^M \mu P_i \parallel \pi\right) \leq \frac{1}{M} \sum_{i=1}^M \text{KL}(\mu P_i \parallel \pi) = \frac{1}{M} \sum_{i=1}^M \text{KL}(\mu^{-i} \parallel \pi^{-i}). \quad (10.2.9)$$

Also, suppose that for each i , ν_i has the same X^{-i} marginal as μ . Then, another application of the chain rule yields

$$\frac{1}{M} \sum_{i=1}^M \text{KL}(\mu^{-i} \parallel \pi^{-i}) \leq \frac{1}{M} \sum_{i=1}^M \text{KL}(\nu_i \parallel \pi). \quad (10.2.10)$$

We pause to provide a motivating remark. Whereas the methods of Section 10.2.1 are understood as sampling analogues of random coordinate descent from optimization, **RS-GS** is the sampling analogue of *alternating minimization*, which successively updates coordinate blocks by exactly minimizing the objective function with respect to those variables. Obviously, alternating minimization possesses an extremal characterization: it produces the largest drop in the function value among all possible updates acting solely on the coordinate block. Similarly, the inequalities (10.2.9) and (10.2.10) are an extremal characterization of **RS-GS**.

Our goal now is to exploit this property by constructing appropriate measures ν_i that witness a drop in the relative entropy. Since we have the marginal constraint $\nu_i^{-i} = \mu^{-i}$, a natural idea is to define the conditional $\nu_i^{i|i-i}$ by “transporting” $\mu^{i|i-i}$ toward $\pi^{i|i-i}$.

Introduce the **Knothe–Rosenblatt map**, defined as follows. First, let T^1 be the optimal transport map from μ^1 to π^1 . Next, for each $1 \leq i < M$, let $T^{i+1}(x^{i+1} \mid x^{1:i})$ be the optimal transport map from $\mu^{i+1|i}(\cdot \mid x^{1:i})$ to $\pi^{i+1|i}(\cdot \mid y^{1:i})$ evaluated at x^{i+1} , where $y^{1:i} := (T^1(x^1), \dots, T^i(x^i \mid x^{1:i-1}))$. By construction, $T : x \mapsto (T^1(x^1), \dots, T^M(x^M \mid x^{1:M-1}))$ is a valid optimal transport map from μ to π (Exercise 10.9). We extend each map T^i to a map $\tilde{T}^i : \mathbb{R}^d \rightarrow \mathbb{R}^d$ via $\tilde{T}^i(x) := (x^{-i}, T^i(x^i \mid x^{1:i-1}))$. Also, let $T_\lambda := (1 - \lambda) \text{id} + \lambda T$ and $\tilde{T}_\lambda^i := (1 - \lambda) \text{id} + \lambda \tilde{T}^i$ for $\lambda \in [0, 1]$. Then, we take the choice $\nu_i := (\tilde{T}_\lambda^i)_\# \mu$. Observe that when $\lambda = 1$, we simply have $\nu_i = \mu P_i$ and (10.2.10) becomes an equality, but this does not make any progress toward understanding **RS-GS**. Instead, we will show that choosing an appropriate value of $\lambda < 1$ leads to the entropy contraction bound (10.2.8).

Decompose the KL divergence according to the energy and the entropy:

$$\text{KL}(\nu \parallel \pi) = \mathcal{E}(\nu) + \mathcal{H}(\nu) := \int \left(V_0(x) + \sum_{i=1}^M V_i(x^i) \right) \nu(dx) + \int \log \nu(x) \nu(dx).$$

Recall that we can absorb the normalizing constant into the definition of V_0 to make this hold true. We will handle the two terms separately.

Energy. By the anisotropic smoothness and convexity assumptions,

$$\begin{aligned} \frac{1}{M} \sum_{i=1}^M V_0(\tilde{T}_\lambda^i(x)) &\leq \frac{1}{M} \sum_{i=1}^M \left[V_0(x) + \lambda \langle \nabla_i V_0(x), T^i(x^i \mid x^{1:i-1}) - x^i \rangle + \frac{\beta_i \lambda^2}{2} \|T^i(x^i \mid x^{1:i-1}) - x^i\|^2 \right] \\ &= \frac{M-1}{M} V_0(x) + \frac{1}{M} \left[V_0(x) + \lambda \langle \nabla V_0(x), T(x) - x \rangle + \frac{\lambda^2}{2} \|T(x) - x\|_\beta^2 \right] \\ &\leq \frac{M-1}{M} V_0(x) + \frac{1}{M} \left[V_0(T_\lambda(x)) + \frac{\lambda^2 (1 - \kappa^{-1})}{2} \|T(x) - x\|_\beta^2 \right]. \end{aligned}$$

For the separable part,

$$\frac{1}{M} \sum_{i=1}^M V_j([\tilde{T}_\lambda^i(x)]^j) = \frac{M-1}{M} V_j(x^j) + \frac{1}{M} V_j((1 - \lambda) x^j + \lambda T^j(x^j \mid x^{1:j-1})),$$

and summing this over $j \in [M]$ yields

$$\frac{1}{M} \sum_{i,j=1}^M V_j([\bar{T}_\lambda^i(x)]^j) = \frac{M-1}{M} \sum_{j=1}^M V_j(x^j) + \frac{1}{M} \sum_{j=1}^M V_j([T_\lambda(x)]^j).$$

Adding these inequalities together and integrating w.r.t. $\mu(dx)$ yields

$$\frac{1}{M} \sum_{i=1}^M \mathcal{E}(v_i) \leq \frac{M-1}{M} \mathcal{E}(\mu) + \frac{1}{M} \left[\mathcal{E}((T_\lambda)_\# \mu) + \frac{\lambda^2 (1 - \kappa^{-1})}{2} \int \|T(x) - x\|_\beta^2 \mu(dx) \right].$$

Moreover, convexity yields

$$\mathcal{E}((T_\lambda)_\# \mu) \leq (1 - \lambda) \mathcal{E}(\mu) + \lambda \mathcal{E}(\pi) - \frac{\lambda (1 - \lambda) \kappa^{-1}}{2} \int \|T(x) - x\|_\beta^2 \mu(dx).$$

Choosing $\lambda = \kappa^{-1}$ produces a cancellation, thus

$$\frac{1}{M} \sum_{i=1}^M \mathcal{E}(v_i) \leq \left(1 - \frac{1}{\kappa M}\right) \mathcal{E}(\mu) + \frac{1}{\kappa M} \mathcal{E}(\pi). \quad (10.2.11)$$

Entropy. Here, we use the fact that T_λ has a triangular Jacobian. By the change of variables formula—see (1.4.3)—we have

$$\begin{aligned} \frac{1}{M} \sum_{i=1}^M \mathcal{H}(v_i) &= \mathcal{H}(\mu) - \frac{1}{M} \sum_{i=1}^M \int \log \det \nabla \bar{T}_\lambda^i d\mu = \mathcal{H}(\mu) - \frac{1}{M} \int \log \det \nabla T_\lambda d\mu \\ &= \frac{M-1}{M} \mathcal{H}(\mu) + \frac{1}{M} \mathcal{H}((T_\lambda)_\# \mu). \end{aligned}$$

Next, we apply convexity of the entropy, which implies $\mathcal{H}((T_\lambda)_\# \mu) \leq (1 - \lambda) \mathcal{H}(\mu) + \lambda \mathcal{H}(\pi)$. Actually, this does not directly follow from (1.4.4), since T is not the optimal transport map; $\lambda \mapsto (T_\lambda)_\# \mu$ is not a Wasserstein geodesic. Nevertheless, the proof of (1.4.4) still goes through in this case, using the fact that ∇T is triangular with positive diagonal entries. See also Exercise 10.17. For $\lambda = \kappa^{-1}$, this yields

$$\frac{1}{M} \sum_{i=1}^M \mathcal{H}(v_i) \leq \left(1 - \frac{1}{\kappa M}\right) \mathcal{H}(\mu) + \frac{1}{\kappa M} \mathcal{H}(\pi). \quad (10.2.12)$$

Finishing the proof. The proof of (10.2.8) now follows from (10.2.9), (10.2.10), (10.2.11), and (10.2.12). $\square$

We also note that Theorem 10.2.7 implies a mixing time bound for the hit-and-run sampler, see Exercise 10.10.

10.3 Mirror Langevin

The *mirror descent* method in optimization changes the geometry of the algorithm via the use of a **mirror map** $\phi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\infty\}$. Here, ϕ is a convex function and we denote $\mathcal{X} := \text{int dom } \phi$; we assume that $\mathcal{X} \neq \emptyset$, that ϕ is strictly convex and differentiable on $\mathcal{X}$, and that ϕ is a barrier for $\mathcal{X}$ in the sense that $\|\nabla \phi(x_n)\| \rightarrow \infty$ whenever $(x_n)_{n \in \mathbb{N}} \subseteq \mathcal{X}$ converges to a point on $\partial \mathcal{X}$. Then, rather than following the gradient descent iteration

$$x_{n+1} := x_n - h \nabla V(x_n), \quad n = 0, 1, 2, \dots \quad (10.3.1)$$

we can instead consider the **mirror descent** iteration

$$\nabla\phi(x_{n+1}) := \nabla\phi(x_n) - h \nabla V(x_n), \quad n = 0, 1, 2, \dots \quad (10.3.2)$$

The assumptions on the mirror map ϕ ensure that the iteration (10.3.2) is well-defined. When $\phi = \frac{\|\cdot\|^2}{2}$, then mirror descent coincides with gradient descent.

Historically, mirror descent was introduced in [Nemirovsky and Yudin \(1983\)](#) with the following intuition. Suppose that we are optimizing a function V which is not defined over the Euclidean space $\mathbb{R}^d$, but rather over a Banach space $\mathcal{B}$. Then, the gradient of V is not an element of $\mathcal{B}$ but rather of the dual space $\mathcal{B}^*$, and so the gradient descent iteration (10.3.1) does not even make sense. On the other hand, the mirror descent iteration (10.3.2) works because the primal point $x_n \in \mathcal{B}$ is first mapped to the dual space $\mathcal{B}^*$ via the mapping $\nabla\phi$. This reasoning is not so esoteric as it may seem, because even for a function V defined over $\mathbb{R}^d$ its natural geometry may correspond to a different norm (e.g., the ℓ_1 norm), in which case $\mathbb{R}^d$ is better viewed as a Banach space.

Our aim is to understand the sampling analogue of mirror descent, known as *mirror Langevin*. We will keep in mind the key example of **constrained sampling**. Here, the potential $V : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\infty\}$ has domain $\mathcal{X} \subseteq \mathbb{R}^d$. In this case, the standard Langevin algorithm leaves the constraint set $\mathcal{X}$ which is undesirable; in particular, it is not possible to obtain guarantees in metrics such as KL divergence because the law of the iterate of the algorithm is not absolutely continuous with respect to the target. Projecting the iterates onto $\mathcal{X}$ does not solve this issue because the law of the iterate will then have positive mass on the boundary $\partial\mathcal{X}$. Besides, Euclidean projection may not adapt well to the shape of the constraint set $\mathcal{X}$. Instead, the use of a mirror map ϕ which is a barrier for $\mathcal{X}$ can *automatically* enforce the constraint.

10.3.1 Continuous-Time Considerations

In continuous time, the mirror Langevin diffusion $(Z_t)_{t \geq 0}$ is the solution to the stochastic differential equation

$$Z_t^* = \nabla\phi(Z_t), \quad dZ_t^* = -\nabla V(Z_t) dt + \sqrt{2} [\nabla^2\phi(Z_t)]^{1/2} dB_t. \quad (10.3.3)$$

Here, the diffusion term is no longer an isotropic Brownian motion but rather involves the matrix $[\nabla^2\phi(Z_t)]^{1/2}$; this is necessary in order to ensure that the stationary distribution is π . Also, we have given the SDE in the dual space. Using Itô's formula ([Theorem 1.1.19](#)), one can write down an SDE for $(Z_t)_{t \geq 0}$ in the primal space, but it is more complicated, involving the third derivative tensor of ϕ (see [Exercise 10.12](#)), and as such we prefer to work with the representation (10.3.3).

Using (10.3.3), we can compute the generator $\mathcal{L}$, the carré du champ Γ , and the Dirichlet energy $\mathcal{E}$ of the mirror Langevin diffusion (see [Section 1.2](#) and [Exercise 10.13](#)):

$$\Gamma(f, g) = \langle \nabla f, [\nabla^2\phi]^{-1} \nabla g \rangle, \quad \mathcal{E}(f, g) = \int \langle \nabla f, [\nabla^2\phi]^{-1} \nabla g \rangle d\pi. \quad (10.3.4)$$

The expression shows that the mirror Langevin diffusion is reversible with respect to π . Also, if π_t denotes the law of Z_t , then

$$\int f \partial_t \pi_t = \partial_t \int f d\pi_t = \int \mathcal{L} f \frac{\pi_t}{\pi} d\pi = - \int \langle \nabla f, [\nabla^2\phi]^{-1} \nabla \frac{\pi_t}{\pi} \rangle d\pi \quad (10.3.5)$$

$$= - \int \langle \nabla f, [\nabla^2\phi]^{-1} \nabla \log \frac{\pi_t}{\pi} \rangle d\pi_t = \int f \operatorname{div}(\pi_t [\nabla^2\phi]^{-1} \nabla \log \frac{\pi_t}{\pi}) \quad (10.3.6)$$

from which we deduce the Fokker–Planck equation

$$\partial_t \pi_t = \operatorname{div}(\pi_t [\nabla^2 \phi]^{-1} \nabla \log \frac{\pi_t}{\pi}).$$

From the interpretation as a continuity equation (see [Theorem 1.3.17](#)), we see that $(\pi_t)_{t \geq 0}$ describes the evolution of a particle which travels according to the family of vector fields $t \mapsto -[\nabla^2 \phi]^{-1} \nabla \log(\pi_t/\pi)$. Recalling that $\nabla \log(\pi_t/\pi)$ is the Wasserstein gradient of $\operatorname{KL}(\cdot \parallel \pi)$ at π_t , we can interpret the mirror Langevin diffusion as a “mirror flow” of the KL divergence in Wasserstein space.

Alternatively, we can equip $\mathcal{X}$ with the Riemannian metric induced by $\nabla^2 \phi$, i.e., we set $\langle u, v \rangle_x := \langle u, \nabla^2 \phi(x) v \rangle$. Then, the mirror Langevin diffusion becomes the Wasserstein gradient flow of the KL divergence over the Riemannian manifold $\mathcal{X}$ (see [Section 2.5.1](#)).

The Newton–Langevin diffusion.

In the special case when the mirror map ϕ is chosen to be the same as the potential V , we arrive at a sampling analogue of Newton’s algorithm, and hence we call it the **Newton–Langevin diffusion**. The equation for the Newton–Langevin diffusion can be written (in the dual space) as

$$dZ_t^* = -Z_t^* dt + \sqrt{2} [\nabla^2 V^*(Z_t^*)]^{-1/2} dB_t, \quad (10.3.7)$$

see [Exercise 10.14](#).

Convergence in continuous time.

In optimization, Newton’s algorithm has many favorable properties. For example, at least locally, it is known that Newton’s algorithm converges *quadratically* rather than linearly, which means that the error at iteration n scales as $\exp(-c_1 \exp(c_2 n))$ for constants $c_1, c_2 > 0$. Also, Newton’s algorithm is *affine-invariant*, meaning that if A is any invertible matrix and we instead apply Newton’s algorithm to the function $\tilde{V}(x) := V(Ax)$, then the iterates $(\tilde{x}_n)_{n \in \mathbb{N}}$ are related to the iterates $(x_n)_{n \in \mathbb{N}}$ of Newton’s algorithm on the original function V via the transformation $\tilde{x}_n = A^{-1}x_n$ ([Exercise 10.15](#)). Consequently, the convergence speed of Newton’s algorithm should not be badly affected by poor conditioning of V .

Can we expect similar properties to hold for the Newton–Langevin diffusion? At least for the property of affine invariance, we have the **Brascamp–Lieb inequality** ([Theorem 2.2.9](#)): if $\pi \propto \exp(-V)$ is strictly log-concave, then for all $f : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\operatorname{var}_\pi f \leq \mathbb{E}_\pi \langle \nabla f, [\nabla^2 V]^{-1} \nabla f \rangle.$$

Below, we will also give an alternative proof of the Bregman transport inequality ([Theorem 2.2.12](#)) based on Wasserstein calculus, which implies the Brascamp–Lieb inequality. Note that in the strongly convex case $\nabla^2 V \geq \alpha I_d$, it implies a Poincaré inequality (in the sense of [Example 1.2.23](#)) for π with constant $1/\alpha$. However, in our present context with Dirichlet energy given by (10.3.4), we instead interpret the Brascamp–Lieb inequality as a Poincaré inequality (in the sense of [Definition 1.2.19](#)) for the Newton–Langevin diffusion. Then, the Poincaré constant is 1, which is notably independent of the strong convexity parameter of V .

We also obtain a Poincaré inequality for the mirror Langevin diffusion under the condition of relative strong convexity.

Definition 10.3.8. Let $\phi, V : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\infty\}$ be convex functions, and assume that $\mathcal{X} = \text{int dom } \phi = \text{int dom } V$. Then:

1 V is α -**relatively convex** (w.r.t. ϕ) if for all $x \in \mathcal{X}$,

$$\nabla^2 V(x) \geq \alpha \nabla^2 \phi(x) .$$

2 V is β -**relatively smooth** (w.r.t. ϕ) if for all $x \in \mathcal{X}$,

$$\nabla^2 V(x) \leq \beta \nabla^2 \phi(x) .$$

Observe that when $\phi = \frac{\|\cdot\|^2}{2}$, these definitions reduce to the usual definitions of strong convexity and smoothness. Recall from [Definition 2.2.11](#) that the Bregman divergence D_ϕ associated with ϕ is the mapping $D_\phi(\cdot, \cdot) : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ given by

$$D_\phi(x, y) := \phi(x) - \phi(y) - \langle \nabla \phi(y), x - y \rangle .$$

The Bregman divergence plays an important role in the analysis of mirror Langevin because it is the correct substitute for the Euclidean distance $(x, y) \mapsto \frac{1}{2} \|x - y\|^2$ in this context. Note the following observations: (1) D_ϕ is non-negative due to convexity of ϕ , and if ϕ is strictly convex then it equals 0 if and only if its two arguments are equal; (2) since $D_\phi(x, y)$ is defined by subtracting the first-order Taylor expansion of ϕ at y , it behaves infinitesimally like a squared distance; in particular,

$$D_\phi(x, y) \sim \frac{1}{2} \langle y - x, \nabla^2 \phi(x) (y - x) \rangle \quad \text{as } y \rightarrow x ;$$

(3) when $\phi = \frac{\|\cdot\|^2}{2}$, then D_ϕ is precisely one-half times the squared Euclidean distance.

Using this definition, we have the following reformulations of relative convexity and relative smoothness ([Exercise 10.16](#)).

Lemma 10.3.9. V is α -relatively convex w.r.t. ϕ if and only if

$$D_V \geq \alpha D_\phi .$$

Similarly, V is β -relatively smooth w.r.t. ϕ if and only if

$$D_V \leq \beta D_\phi .$$

Returning to the mirror Langevin diffusion, the following corollary is an immediate consequence of the Brascamp–Lieb inequality and the definition of relative convexity.

Corollary 10.3.10 (Mirror Poincaré inequality, [Chewi et al. \(2020\)](#)). Suppose that the potential V is α -relatively convex w.r.t. ϕ . Then, the mirror Langevin diffusion satisfies the following Poincaré inequality: for all $f : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\text{var}_\pi f \leq \frac{1}{\alpha} \mathbb{E}_\pi \langle \nabla f, [\nabla^2 \phi]^{-1} \nabla f \rangle = \frac{1}{\alpha} \mathcal{E}(f, f) .$$

So far so good: we have defined relative convexity, which is a natural generalization of strong convexity and well-studied in the optimization literature; and we have shown that it implies a Poincaré inequality for the mirror Langevin diffusion.

However, the analogies with the standard Langevin diffusion stop here. There are counterexamples which show that relative convexity does *not* imply a log-Sobolev inequality for the mirror Langevin diffusion. This may seem to contradict the Bakry–Émery theorem (Theorem 1.2.30), which holds for any Markov diffusion. The issue here is that the assumption of relative convexity is *not* a curvature-dimension condition. Indeed, in order to properly formulate a curvature-dimension condition for the Hessian manifold $\mathcal{X}$ equipped with the Riemannian metric $\mathbf{g}$ induced by $\nabla^2\phi$, one must check the $\text{CD}(\alpha, \infty)$ condition $\nabla^2\tilde{V} + \text{Ric} \geq \alpha$. Here, ∇^2 denotes the Riemannian Hessian and $\tilde{V}$ is the potential corresponding to the relative density of π w.r.t. the Riemannian volume measure on $(\mathcal{X}, \mathbf{g})$; see Section 2.5. Such a calculation was performed in, e.g., Kolesnikov (2014), which implies that if $(\nabla\phi)_\# \pi$ is log-concave, then the Newton–Langevin diffusion satisfies $\text{CD}(\frac{1}{2}, \infty)$. However, it is not clear under what conditions $(\nabla\phi)_\# \pi$ is log-concave.

Here is another consequence of the fact that relative convexity is not a curvature-dimension condition: relative convexity (apparently) does not seem to imply contraction properties for the mirror Langevin diffusion with respect to an appropriately defined Wasserstein metric.

To summarize: either we can assume the curvature-dimension condition $\text{CD}(\alpha, \infty)$, which imposes complicated conditions on ϕ and V , or we can adopt the more interpretable relative convexity assumption, which in turn only implies a Poincaré inequality (Corollary 10.3.10). We will follow the latter approach.

Upon reflection, the curvature-dimension approach for studying the mirror Langevin diffusion is arguably the less natural one. Indeed, the curvature-dimension approach is based on viewing the mirror Langevin diffusion from the lens of Riemannian geometry, but the mirror descent algorithm in optimization is not typically studied via Riemannian geometry. Instead, the study of mirror descent is based on ideas from convex analysis, centered around the Bregman divergence. So it seems prudent at this stage to abandon the Riemannian interpretation of mirror Langevin in favor of convex analysis tools, and this is indeed how our discretization proof will go. In fact, in lieu of using the Poincaré inequality in Corollary 10.3.10, we will *directly* use relative convexity.

10.3.2 Discretization Preliminaries

Following Ahn and Chewi (2021), we consider the following discretization of (10.3.3).

$$\begin{aligned} \nabla\phi(X_{nh}^+) &:= \nabla\phi(X_{nh}) - h \nabla V(X_{nh}), \\ X_{(n+1)h} &:= \nabla\phi^*(X_{(n+1)h}^*), \end{aligned} \tag{MLMC}$$

where

$$X_t^* = \nabla\phi(X_{nh}^+) + \sqrt{2} \int_{nh}^t [\nabla^2\phi^*(X_s^*)]^{-1/2} dB_s \quad \text{for } t \in [nh, (n+1)h]. \tag{10.3.11}$$

Note that when $\phi = \frac{\|\cdot\|^2}{2}$, this reduces to the standard LMC algorithm. When generalizing LMC to different mirror maps, this discretization is chosen to preserve the “forward-flow” interpretation of Wibisono (2018) (see Section 4.3). In particular, the update from X_{nh} to X_{nh}^+ is a mirror descent step, while the update from X_{nh}^+ to $X_{(n+1)h}$ follows a “Wasserstein mirror flow” of the (negative) entropy.

However, implementing MLMC requires the exact simulation of a diffusion process, so is it truly a “discretization”? To address this, note that simulating the mirror diffusion (10.3.11) does not require additional queries to the potential V , since it only depends on the mirror map ϕ (except for the initialization). To an extent, any algorithm based on mirror maps requires the implementation of

certain primitive operations involving ϕ , such as computation of $\nabla\phi$ or inversion of $\nabla\phi$; in practice, this requires ϕ to have a “simple” structure such that these operations have closed-form expressions, or are at least cheap enough to be negligible relative to the cost of computing the gradient of V . In our consideration of **MLMC**, we take the diffusion (10.3.11) to be another primitive operation associated with the mirror map ϕ . This is indeed appropriate for many applications, e.g., when ϕ is a separable function $\phi : x \mapsto \sum_{i=1}^d \phi_i(x_i)$, and it will streamline our technical analysis. Nevertheless, implementation of **MLMC** remains a key obstacle to its practicality and necessitates further research in this direction.

Our approach to studying **MLMC** is to adapt the convex optimization approach introduced in Section 4.3.

Key technical results.

We now establish the analogues of the various facts that we invoked in the study of **LMC**.

First, in the standard gradient descent analysis, if $x_+ := x - h \nabla V(x)$, then we have the key inequality

$$\langle \nabla V(x), x^+ - z \rangle = \frac{1}{2h} \{ \|x - z\|^2 - \|x^+ - z\|^2 - \|x^+ - x\|^2 \} \quad \text{for all } z \in \mathbb{R}^d.$$

Remarkably, there is an analogue of this fact for mirror descent, which follows from the following identity (which can be checked by simple algebra, see [Exercise 10.16](#)):

$$\langle \nabla\phi(\tilde{x}) - \nabla\phi(x), \tilde{x} - z \rangle = D_\phi(\tilde{x}, x) + D_\phi(z, \tilde{x}) - D_\phi(z, x), \quad \text{for all } x, \tilde{x}, z \in \mathcal{X}. \quad (10.3.12)$$

Lemma 10.3.13 (Bregman proximal lemma, [Chen and Teboulle \(1993\)](#)). *For $x \in \mathcal{X}$, let x^+ be defined via $\nabla\phi(x^+) = \nabla\phi(x) - h \nabla V(x)$. Then, for all $z \in \mathcal{X}$,*

$$\langle \nabla V(x), x^+ - z \rangle = \frac{1}{h} \{ D_\phi(z, x) - D_\phi(z, x^+) - D_\phi(x^+, x) \}.$$

Proof Note that

$$\langle \nabla V(x), x^+ - z \rangle = -\frac{1}{h} \langle \nabla\phi(x^+) - \nabla\phi(x), x^+ - z \rangle.$$

Substituting the identity (10.3.12) into the above equation proves the result. $\square$

Unlike the case of **LMC**, the presence of a non-constant diffusion matrix involving $\nabla^2\phi$ introduces another source of discretization error. To address this, we introduce a condition on the third derivative of ϕ . Note however that a uniform bound on the operator norm of $\nabla^3\phi$ is not compatible with the assumption that ϕ tends to $+\infty$ on $\partial\mathcal{X}$. The solution to this issue was discovered in [Nesterov and Nemirovskii \(1994\)](#): we can ask that $\nabla^3\phi$ is bounded *with respect to the geometry induced by ϕ* . This approach also has the benefit of being consistent with the affine invariance of Newton’s method. The precise definition of the third derivative condition is as follows.

Definition 10.3.14 (Self-concordance). The mirror map ϕ is said to be M_ϕ -**self-concordant** if for all $x \in \mathcal{X}$ and all $v \in \mathbb{R}^d$,

$$\nabla^3\phi(x)[v, v, v] \leq 2M_\phi \|v\|_{\nabla^2\phi(x)}^3 := 2M_\phi \langle v, \nabla^2\phi(x) v \rangle^{3/2}.$$

The norm $\|v\|_{\nabla^2 \phi(x)} := \sqrt{\langle v, \nabla^2 \phi(x) v \rangle}$ is called the *local norm*, and it is the tangent space norm for the Riemannian metric induced by ϕ .

The definition implies the following result, stated without proof:⁵

Lemma 10.3.15 (Nesterov (2018, Corollary 5.1.1)). *Suppose that ϕ is M_ϕ -self-concordant. Then, for all $x \in \mathcal{X}$ and $u \in \mathbb{R}^d$,*

$$\nabla^3 \phi(x) u \leq 2M_\phi \|u\|_{\nabla^2 \phi(x)} \nabla^2 \phi(x) .$$

Self-concordant functions are well-studied due to their central role in the theory of interior-point methods for optimization, see the monograph Nesterov and Nemirovskii (1994). A key example of a self-concordant mirror map is when the constraint set is a polytope,

$$\mathcal{X} = \{x \in \mathbb{R}^d : \langle a_i, x \rangle < b_i \text{ for all } i \in [N]\} ,$$

in which case $\phi(x) = \log(1/\sum_{i=1}^N (b_i - \langle a_i, x \rangle))$ is self-concordant with $M_\phi = 1$.

Finally, a key step in our analysis of LMC was to use the convexity of the entropy functional along W_2 geodesics. In our analysis of MLMC, we will replace the W_2 distance with the Bregman transport cost $\mathcal{D}_V$ (recall Definition 2.2.11). To study these costs, we first state an analogue of Brenier's theorem (Theorem 1.3.8).

Theorem 10.3.16 (Brenier's theorem for the Bregman transport cost). *Suppose that $\mu, \nu \in \mathcal{P}(\mathbb{R}^d)$. Then, the unique optimal Bregman transport coupling (X, Y) for μ and ν is of the form*

$$\nabla \phi(Y) = \nabla \phi(X) - \nabla h(X) ,$$

where $h : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{-\infty\}$ is such that $\phi - h$ is convex.

Proof sketch We need facts about optimal transport with general costs c (Exercise 1.10). Namely, the optimal pair of dual potentials (f^*, g^*) are c -conjugates, meaning that

$$\begin{aligned} f^*(x) &= \inf_{y \in \mathbb{R}^d} \{c(x, y) - g^*(y)\} , \\ g^*(y) &= \inf_{x \in \mathbb{R}^d} \{c(x, y) - f^*(x)\} . \end{aligned}$$

For γ^* -a.e. (x, y) , it holds that $f^*(x) + g^*(y) = c(x, y)$. If c is smooth and such that $\nabla_x c(x, \cdot)$ is injective for all $x \in \mathbb{R}^d$, then from the definition of g^* it suggests that we have $\nabla_x c(x, y) = \nabla f^*(x)$ for γ^* -a.e. (x, y) . See Villani (2009b, Theorem 10.28) for a rigorous statement and proof of these results.

Applying this to our cost function $c = D_\phi$, it yields the existence of D_ϕ -conjugates $h, \tilde{h} : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{-\infty\}$ such that $\nabla_x D_\phi(x, y) = \nabla h(x)$ under the optimal plan γ^* . Hence,

$$\nabla \phi(y) = \nabla \phi(x) - \nabla h(x) , \quad \text{for } \gamma^*\text{-a.e. } (x, y)$$

and

$$h(x) = \inf_{y \in \mathbb{R}^d} \{D_\phi(x, y) - \tilde{h}(y)\} .$$

⁵The proof is surprisingly difficult. The reader can try to prove the result with a worse constant factor.

By expanding out the definition of D_ϕ , we rewrite this as

$$\phi(x) - h(x) = \sup_{y \in \mathcal{X}} \{ \langle \nabla \phi(y), x - y \rangle + \tilde{h}(y) + \phi(y) \}.$$

As a supremum of affine functions, it follows that $\phi - h$ is convex. $\square$

Recall that the usual W_2 geodesic from μ_0 to μ_1 is given as follows: there is a convex function $\varphi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{\infty\}$ such that for $X_0 \sim \mu_0$, the pair $(X_0, X_1) := (X_0, \nabla \varphi(X_0))$ is an optimal coupling. By taking the linear interpolation $X_t := (1 - t) X_0 + t X_1$ and setting $\mu_t := \text{law}(X_t)$, we obtain the W_2 geodesic $(\mu_t)_{t \in [0,1]}$.

Now consider the above theorem. If we let $W := \nabla \phi(Y) = \nabla \phi(X) - \nabla h(X)$, then we have the coupling $(X_0, X_1) := (\nabla(\phi - h)^*(W), \nabla \phi^*(W))$ of $\mu_0 := \mu$ and $\mu_1 := \nu$. Then, we can interpolate by setting $X_t := (1 - t) X_0 + t X_1$ and $\mu_t := \text{law}(X_t)$, which defines an alternative path joining μ_0 to μ_1 . Note that (X_0, W) and (X_1, W) are both optimally coupled for the W_2 distance (since $\phi - h$ is convex). This is a special case of the following.

Definition 10.3.17. A curve $(\mu_t)_{t \in [0,1]} \subseteq \mathcal{P}_2(\mathbb{R}^d)$ is a **generalized geodesic** if there exists another measure $\rho \in \mathcal{P}_2(\mathbb{R}^d)$ such that

$$\mu_t = \text{law}(X_t), \quad X_t := (1 - t) X_0 + t X_1,$$

and (X_0, W) , (X_1, W) are both optimally coupled for the W_2 metric, where $W \sim \rho$.

It is left as [Exercise 10.17](#) to check that the entropy functional is also convex along generalized geodesics.

Theorem 10.3.18. Let $\mathcal{H}(\mu) := \int \mu \log \mu$. Then, for any generalized geodesic $(\mu_t)_{t \in [0,1]}$,

$$t \mapsto \mathcal{H}(\mu_t) \quad \text{is convex}.$$

In our context, it implies that if $\mu, \nu \in \mathcal{P}(\mathbb{R}^d)$ and (X, Y) is an optimal coupling for the Bregman cost $\mathcal{D}_\phi(\mu, \nu)$, then

$$\mathcal{H}(\nu) \geq \mathcal{H}(\mu) + \mathbb{E} \langle \nabla \log \mu(X), Y - X \rangle. \quad (10.3.19)$$

As an aside, we remark that this observation leads to another proof of the Bregman transport inequality.

Proof of the Bregman transport inequality (Theorem 2.2.12) Let $\pi = \exp(-V)$, where V is strictly convex, and let $\mu \in \mathcal{P}(\mathbb{R}^d)$. Let $X \sim \mu$, $Z \sim \pi$ be optimally coupled for the Bregman transport cost. Then,

$$\text{KL}(\mu \parallel \pi) = \mathbb{E} V(X) + \mathcal{H}(\mu).$$

For the first term, by the definition of D_V ,

$$\mathbb{E} V(X) = \mathbb{E} V(Z) + \mathbb{E} D_V(X, Z) + \mathbb{E} \langle \nabla V(Z), X - Z \rangle.$$

For the second term, (10.3.19) implies

$$\mathcal{H}(\mu) \geq \mathcal{H}(\pi) + \mathbb{E} \langle \nabla \log \pi(Z), X - Z \rangle.$$

Hence,

$$\begin{aligned} \text{KL}(\mu \parallel \pi) &\geq \underbrace{\mathbb{E} V(Z) + \mathcal{H}(\pi)}_{=\text{KL}(\pi \parallel \pi)=0} + \underbrace{\mathbb{E} \langle \nabla V(Z) + \nabla \log \pi(Z), X - Z \rangle}_{=0} + \mathbb{E} D_V(X, Z) \\ &= \mathcal{D}_V(\mu \parallel \pi), \end{aligned}$$

which is what we wanted to show. $\square$

10.3.3 Discretization Analysis

Analysis for the smooth case.

We now prove the following result.

Theorem 10.3.20 (Ahn and Chewi (2021)). Suppose that $\pi = \exp(-V)$ is the target distribution and that ϕ is the mirror map. Assume:

- V is α -relatively convex and β -relatively smooth w.r.t. ϕ .
- ϕ is M_ϕ -self-concordant.
- V is L -relatively Lipschitz w.r.t. ϕ , i.e., $\|\nabla V(x)\|_{[\nabla^2 \phi(x)]^{-1}} \leq L$ for all $x \in \mathcal{X}$.

Let $(\mu_{nh})_{n \in \mathbb{N}}$ denote the law of MLMC and let $\beta' := \beta + 2LM_\phi$.

1 (Weakly convex case) Suppose that $\alpha = 0$. For any $\varepsilon \in [0, \sqrt{d}]$, if we take step size $h \asymp \varepsilon^2/\beta'd$, then for the mixture distribution $\bar{\mu}_{Nh} := N^{-1} \sum_{n=1}^N \mu_{nh}$ it holds that $\sqrt{\text{KL}(\bar{\mu}_{Nh} \parallel \pi)} \leq \varepsilon$ after

$$N = O\left(\frac{\beta'd \mathcal{D}_\phi(\pi, \mu_0)}{\varepsilon^4}\right) \quad \text{iterations}.$$

2 (Strongly convex case) Suppose that $\alpha > 0$ and let $\kappa := \beta'/\alpha$ denote the “condition number”. Then, for any $\varepsilon \in [0, \sqrt{d}]$, with step size $h \asymp \alpha \varepsilon^2/\beta'd$ we obtain $\sqrt{\alpha} \mathcal{D}_\phi(\pi, \mu_{Nh}) \leq \varepsilon$ and $\sqrt{\text{KL}(\bar{\mu}_{Nh, 2Nh} \parallel \pi)} \leq \varepsilon$ after

$$N = O\left(\frac{\kappa d}{\varepsilon^2} \log \frac{\alpha \mathcal{D}_\phi(\pi, \mu_0)}{\varepsilon^2}\right) \quad \text{iterations},$$

where $\bar{\mu}_{Nh, 2Nh} := N^{-1} \sum_{n=N+1}^{2N} \mu_{nh}$.

Similarly to Theorem 4.3.6, the theorem follows from a key recursion.

Lemma 10.3.21. Under the assumptions of Theorem 10.3.20, if $h \leq 1/\beta$, then

$$h \text{KL}(\mu_{(n+1)h} \parallel \pi) \leq (1 - \alpha h) \mathcal{D}_\phi(\pi, \mu_{nh}) - \mathcal{D}_\phi(\pi, \mu_{(n+1)h}) + \beta' dh^2.$$

We proceed to prove the lemma.

Proof We follow the proof of Theorem 4.3.6, indicating the changes necessary to adapt the proof to MLMC. Recall that $\mathcal{E}(\mu) := \int V d\mu$.

1. The forward step dissipates the energy. Let $Z \sim \pi$ be optimally coupled to X_{nh} . Then, applying

the relative convexity and relative smoothness of V ,

$$\begin{aligned} \mathcal{E}(\mu_{nh}^+) - \mathcal{E}(\pi) &= \mathbb{E}[V(X_{nh}^+) - V(X_{nh}) + V(X_{nh}) - V(Z)] \\ &\leq \mathbb{E}[\langle \nabla V(X_{nh}), X_{nh}^+ - X_{nh} \rangle + \beta D_\phi(X_{nh}^+, X_{nh}) \\ &\quad + \langle \nabla V(X_{nh}), X_{nh} - Z \rangle - \alpha D_\phi(Z, X_{nh})] \\ &= \mathbb{E}[\langle \nabla V(X_{nh}), X_{nh}^+ - Z \rangle + \beta D_\phi(X_{nh}^+, X_{nh}) - \alpha D_\phi(Z, X_{nh})]. \end{aligned} \quad (10.3.22)$$

Next, by the Bregman proximal lemma (Lemma 10.3.13),

$$\langle \nabla V(X_{nh}), X_{nh}^+ - Z \rangle = \frac{1}{h} \{D_\phi(Z, X_{nh}) - D_\phi(Z, X_{nh}^+) - D_\phi(X_{nh}^+, X_{nh})\}.$$

Substituting this into (10.3.22) and using $h \leq 1/\beta$, it yields

$$\mathcal{E}(\mu_{nh}^+) - \mathcal{E}(\pi) \leq \frac{1}{h} \{(1 - \alpha h) \mathcal{D}_\phi(\mu_{nh}, \pi) - \mathcal{D}_\phi(\mu_{nh}^+, \pi)\}. \quad (10.3.23)$$

2. The flow step does not substantially increase the energy. We write

$$\mathcal{E}(\mu_{(n+1)h}) - \mathcal{E}(\mu_{nh}^+) = \mathbb{E}[V(\nabla \phi^*(X_{(n+1)h}^*)) - V(\nabla \phi^*(X_{nh}^*))].$$

Let $f(x) := V(\nabla \phi^*(x))$ and apply Itô's formula. Note that

$$\begin{aligned} \nabla f(x) &= \nabla V(\nabla \phi^*(x))^\top \nabla^2 \phi^*(x) = \nabla V(\nabla \phi^*(x))^\top [\nabla^2 \phi(\nabla \phi^*(x))]^{-1}, \\ \nabla^2 f(x) &= [\nabla^2 V(\nabla \phi^*(x))] [\nabla^2 \phi(\nabla \phi^*(x))]^{-1} [\nabla^2 \phi^*(x)] \\ &\quad - \nabla V(\nabla \phi^*(x))^\top [\nabla^2 \phi(\nabla \phi^*(x))]^{-1} [\nabla^3 \phi(\nabla \phi^*(x))] [\nabla^2 \phi(\nabla \phi^*(x))]^{-2}. \end{aligned}$$

Itô's formula decomposes $f(X_{(n+1)h}^*) - f(X_{nh}^*)$ into the sum of an integral and a stochastic integral; since the latter has mean zero, we focus on the first term.

$$\begin{aligned} &\mathbb{E}[f(X_{(n+1)h}^*) - f(X_{nh}^*)] \\ &= \mathbb{E} \int_{nh}^{(n+1)h} \langle \nabla^2 V(X_t) [\nabla^2 \phi(X_t)]^{-2}, \nabla^2 \phi(X_t) \rangle dt \\ &\quad - \mathbb{E} \int_{nh}^{(n+1)h} \langle \nabla V(X_t)^\top [\nabla^2 \phi(X_t)]^{-1} [\nabla^3 \phi(X_t)] [\nabla^2 \phi(X_t)]^{-2}, \nabla^2 \phi(X_t) \rangle dt \\ &= \mathbb{E} \int_{nh}^{(n+1)h} \langle \nabla^2 V(X_t), [\nabla^2 \phi(X_t)]^{-1} \rangle dt \\ &\quad - \mathbb{E} \int_{nh}^{(n+1)h} \text{tr}(\nabla V(X_t)^\top [\nabla^2 \phi(X_t)]^{-1} [\nabla^3 \phi(X_t)] [\nabla^2 \phi(X_t)]^{-1}) dt. \end{aligned} \quad (10.3.24)$$

$$- \mathbb{E} \int_{nh}^{(n+1)h} \text{tr}(\nabla V(X_t)^\top [\nabla^2 \phi(X_t)]^{-1} [\nabla^3 \phi(X_t)] [\nabla^2 \phi(X_t)]^{-1}) dt. \quad (10.3.25)$$

By relative smoothness, since $\nabla^2 V \leq \beta \nabla^2 \phi$,

$$(10.3.24) \leq \beta dh.$$

For (10.3.25), we use Lemma 10.3.15, which implies

$$\begin{aligned} (10.3.25) &\leq 2M_\phi \int_{nh}^{(n+1)h} \mathbb{E}[\|[\nabla^2 \phi(X_t)]^{-1} \nabla V(X_t)\|_{\nabla^2 \phi(X_t)} \text{tr}([\nabla^2 \phi(X_t)] [\nabla^2 \phi(X_t)]^{-1})] dt \\ &\leq 2M_\phi d \int_{nh}^{(n+1)h} \mathbb{E}[\|\nabla V(X_t)\|_{[\nabla^2 \phi(X_t)]^{-1}}] dt \leq 2LM_\phi dh. \end{aligned}$$

Hence, we have proven

$$\mathcal{E}(\mu_{(n+1)h}) - \mathcal{E}(\mu_{nh}^+) \leq \beta' dh. \quad (10.3.26)$$

3. The flow step dissipates the entropy. Let $\mu_t := \text{law}(X_t) = \text{law}(\nabla\phi^*(X_t^*))$. The mirror diffusion (10.3.11) evolves according to the vector field $-[\nabla^2\phi]^{-1} \nabla \log \mu_t$. Also, note that $\nabla_y D_\phi(x, y) = -\nabla^2\phi(x)(y - x)$. Using these, one can show that

$$\partial_t \mathcal{D}_\phi(\pi, \mu_t) \leq \mathbb{E} \langle [\nabla^2\phi(X_t)]^{-1} \nabla \log \mu(X_t), [\nabla^2\phi(X_t)](Z - X_t) \rangle = \mathbb{E} \langle \nabla \log \mu(X_t), (Z - X_t) \rangle,$$

where (Z, X_t) is an optimal coupling for $\mathcal{D}_\phi(\pi, \mu_t)$. Using the convexity of $\mathcal{H}$ along generalized geodesics (Theorem 10.3.18),

$$\mathcal{H}(\pi) - \mathcal{H}(\mu_t) \geq \mathbb{E} \langle \nabla \log \mu(X_t), (Z - X_t) \rangle.$$

Using the fact that $t \mapsto \mathcal{H}(\mu_t)$ is decreasing (prove this from the Fokker–Planck equation for the mirror diffusion!), we then have

$$\mathcal{D}_\phi(\pi, \mu_{(n+1)h}) - \mathcal{D}_\phi(\pi, \mu_{nh}^+) \leq h \{ \mathcal{H}(\pi) - \mathcal{H}(\mu_{(n+1)h}) \}. \quad (10.3.27)$$

Concluding the proof. Combine (10.3.23), (10.3.26), and (10.3.27) to conclude. $\square$

To apply Theorem 10.3.20 to the problem of constrained sampling, we can choose ϕ to be a logarithmic barrier for $\mathcal{X}$, which is self-concordant. In the special case when $V = \phi$, the condition that ϕ is L -relatively Lipschitz with respect to itself is commonly expressed as saying that ϕ is a *self-concordant barrier* with parameters (L, M_ϕ) . Self-concordant barriers are also a core part of the theory of interior point methods; in particular, it is known that every convex body in $\mathbb{R}^d$ admits a $(d, 2)$ -self-concordant barrier, and that this is optimal (see Chewi, 2021; Lee and Yue, 2021). However, this situation is “cheating” because if we want to sample from $\pi \propto \exp(-\phi)$, it does not make sense to assume we can exactly simulate the mirror diffusion associated with ϕ .

Result for the non-smooth case.

Although the preceding result applies when ϕ is a logarithmic barrier, it does not apply to perhaps one of the most classical applications of mirror descent: namely, $\mathcal{X}$ is the probability simplex in $\mathbb{R}^d$ and $\phi(x) := \sum_{i=1}^d x_i \log x_i$ is the entropy. The next result we formulate adopts assumptions which precisely match the usual ones for mirror descent in this context.

Theorem 10.3.28 (Ahn and Chewi (2021)). Suppose that $\pi = \exp(-V)$ is the target distribution and that ϕ is the mirror map. Let $\|\cdot\|$ be a norm on $\mathbb{R}^d$. Assume:

- V is convex and L -Lipschitz w.r.t. the dual norm $\|\cdot\|_*$, in the sense that

$$\|\nabla V(x)\|_* \leq L \quad \text{for all } x \in \mathcal{X}.$$

- ϕ is 1-strongly convex w.r.t. $\|\cdot\|$.

Let $(\mu_{nh})_{n \in \mathbb{N}}$ denote the law of MLMC. For any $\varepsilon > 0$, if we take step size $h \asymp \varepsilon^2/L^2$, then for the mixture $\bar{\mu}_{Nh} := N^{-1} \sum_{n=1}^N \mu_{nh}$ it holds that $\sqrt{\text{KL}(\bar{\mu}_{Nh} \parallel \pi)} \leq \varepsilon$ after

$$N = O\left(\frac{L^2 \mathcal{D}_\phi(\pi, \mu_0^+)}{\varepsilon^4}\right) \quad \text{iterations}.$$

For example, it is a classical fact that the entropy is strongly convex w.r.t. the ℓ_1 norm. We leave the proof of the non-smooth case as [Exercise 10.18](#).

Bibliographical Notes

The complexity of SG-LMC was studied in [Dalalyan \(2017a\)](#); [Dalalyan and Karagulyan \(2019\)](#). The latter paper notes that their W_2 bias bound scales as $O(\sigma^2)$ rather than $O(\sigma)$ for small σ , but this does not seem to imply an asymptotic improvement over (10.1.1). The result of [Exercise 10.1](#) is from [Durmus et al. \(2019\)](#).

The literature on stochastic gradient methods is enormous and the following lists are surely not exhaustive. In the general stochastic gradient setting, some works include [Raginsky et al. \(2017\)](#); [Cheng et al. \(2018\)](#); [Salim et al. \(2019\)](#); [Barkhagen et al. \(2021\)](#); [Zou and Gu \(2021\)](#); [Akyildiz and Sabanis \(2024\)](#); [Lu et al. \(2025\)](#). For the specific finite-sum setting, additional notable works include [Dubey et al. \(2016\)](#); [Chatterji et al. \(2018\)](#); [Zou et al. \(2018\)](#).

The first lower bound for the stochastic gradient setting was shown in [Chatterji et al. \(2022\)](#). In that work, the authors obtained a lower bound on the complexity of sampling using a *stochastic* oracle. Namely, in order to output an approximate sample (to constant error in TV distance) from an $\frac{\alpha}{2}$ -strongly log-concave and α -log-smooth distribution whose mean lies in the ball $B(0, 1/\alpha)$, with an oracle that given $x \in \mathbb{R}^d$ outputs $\nabla V(x) + \xi$ with $\xi \sim \text{normal}(0, \Sigma)$ and $\text{tr } \Sigma \leq \sigma^2 d$, the number of queries required is at least $\Omega(\sigma^2 d)$. This lower bound was revisited in [Chen et al. \(2026a\)](#); see the discussion therein for further comparison with [Theorem 10.1.2](#).

Perhaps the earliest non-asymptotic analysis of coordinate Langevin was [Shen et al. \(2019\)](#). Subsequently, the theory was largely developed in [Ding and Li \(2021\)](#); [Ding et al. \(2021a,b\)](#). Another variant is to precondition by a random subspace at each iteration ([Maunu and Yao, 2025](#)).

Mixing time bounds for the Gibbs sampler were also obtained under isoperimetric assumptions in [Goyal et al. \(2026\)](#).

In the context of optimization, self-concordant barriers play a key role in interior point methods for constrained optimization ([Nesterov and Nemirovskii, 1994](#); [Bubeck, 2015](#); [Nesterov, 2018](#)). Relative convexity and relative smoothness were introduced in [Bauschke et al. \(2017\)](#); [Lu et al. \(2018\)](#).

The first use of mirror maps with the Langevin diffusion was via the *mirrored* Langevin algorithm (which is different from the *mirror* Langevin diffusion) in [Hsieh et al. \(2018\)](#). The mirror Langevin diffusion was introduced in an earlier draft of [Hsieh et al. \(2018\)](#), as well as in [Zhang et al. \(2020\)](#). The work [Zhang et al. \(2020\)](#) also studied the Euler–Maruyama discretization of the mirror Langevin diffusion (which differs from MLMC in that it discretizes the diffusion step as well), but they were unable to prove convergence of the algorithm; they were only able to prove convergence to a Wasserstein ball of non-vanishing radius around π , even as the step size tends to zero. They also conjectured that the non-vanishing bias of the algorithm is unavoidable. Subsequently, [Chewi et al. \(2020\)](#) studied the mirror Langevin diffusion in continuous time, and [Ahn and Chewi \(2021\)](#) introduced and studied the MLMC discretization, which *does* lead to vanishing bias (as $h \searrow 0$).

Since then, there have been further works studying the non-vanishing bias issue: [Jiang \(2021\)](#) studied both the Euler–Maruyama and MLMC discretizations under a “mirror log-Sobolev inequality” and was only able to prove vanishing bias for the latter discretization; [Li et al. \(2022a\)](#) showed that the Euler–Maruyama discretization has vanishing bias under stronger assumptions; [Gatmiry and Vempala \(2022\)](#) studied MLMC as a special case of more general Riemannian Langevin algorithms; and [Daaloul et al. \(2025\)](#) established convergence of Euler–Maruyama under a mirror LSI, albeit with

a slow convergence rate. The bias issue is still not settled, and it is certainly of interest to obtain better guarantees for fully discretized algorithms. Nevertheless, in our presentation, we have stuck with the analysis of [Ahn and Chewi \(2021\)](#) because it is the cleanest, and because it relies on assumptions which are well-motivated from convex optimization.

More generally, the mirror Langevin diffusion is an example of a Riemannian diffusion. See [Hsu \(2002\)](#) for general theory, and [Girolami and Calderhead \(2011\)](#) for applications to Bayesian problems.

Exercises

Stochastic Gradients

▷ Exercise 10.1 (Stochastic gradient analysis via convex optimization)

Extend the analysis of [Theorem 4.3.6](#) to handle stochastic gradients, and compare the obtained result with (10.1.1).

▷ Exercise 10.2 (Lower bound for stochastic gradients)

Here, we prove [Theorem 10.1.2](#). Let $p \in [0, 1]$ be such that $p\psi(\varepsilon/p) + \psi(\varepsilon) \leq 1$ and consider the following stochastic gradient oracles:

- For $\text{normal}(0, 1)$, given $x \in \mathbb{R}$, the oracle returns x .
- For $\text{normal}(\varepsilon, 1)$, given $x \in \mathbb{R}$, the oracle returns $x - \frac{\varepsilon}{p}$ with probability p , and x otherwise.

- 1 Show that these stochastic gradient oracles satisfy the assumptions of [Theorem 10.1.2](#).
- 2 Let $\hat{\mu}_0, \hat{\mu}_\varepsilon$ denote the output distributions of the algorithm after N queries to the two oracles respectively. Show that $\|\hat{\mu}_0 - \hat{\mu}_\varepsilon\|_{\text{TV}} \leq pN$.
- 3 Give a lower bound on $\|\text{normal}(0, 1) - \text{normal}(\varepsilon, 1)\|_{\text{TV}}$. Then, use the triangle inequality to prove [Theorem 10.1.2](#).

▷ Exercise 10.3 (Unbounded variance for the naïve gradient estimator)

Consider the naïve stochastic gradient estimator in the finite-sum setting of [Section 10.1.2](#) with $B = 1$. Construct $\{f_i\}_{i=1}^m$ such that the variance at stationarity, $m^{-1} \sum_{i=1}^m \mathbb{E}_\pi [\|\nabla f_i - \nabla f\|^2]$, cannot be bounded in terms of the problem parameters α, β, d , and m alone.

▷ Exercise 10.4 (Control variate LMC)

Use [Lemma 5.1.2](#) to establish (10.1.4).

Coordinate Methods

▷ Exercise 10.5 (Analysis of CG-LMC)

Use [Lemma 5.1.2](#) to establish (10.2.1).

▷ Exercise 10.6 (Analysis of RC-LMC)

Use the strategy outlined in the main text together with [Lemma 5.1.2](#) to establish (10.2.3).

Hint: To leading order in h , $\tilde{\mathcal{E}}_{\text{weak}} \lesssim h^{3/2} \sqrt{\sum_{i=1}^d \beta_i^2 / p_i}$ and $\tilde{\mathcal{E}}_{\text{strong}} \lesssim h^{3/2} \sqrt{\sum_{i=1}^d \beta_i^2 / p_i^2}$.

▷ Exercise 10.7 (Analysis of SVRG-LMC)

Rigorously establish (10.2.4).

▷ **Exercise 10.8** (SVRG–ULMC)

Formulate an algorithm which combines SVRG with ULMC and establish the bound (10.2.5).

▷ **Exercise 10.9** (Knothe–Rosenblatt map)

Check that the Knothe–Rosenblatt map, defined in the proof of Theorem 10.2.7, is a valid transport map from μ to π .

▷ **Exercise 10.10** (Hit-and-run)

Consider the hit-and-run sampler. At each iteration, it picks a random line $\{X + \lambda V : \lambda \in \mathbb{R}\}$ passing through the current point X , where V is a unit vector chosen uniformly at random. Then, the next iterate is $X + \Lambda V$, where Λ is drawn according to the distribution with density proportional to $\lambda \mapsto \pi(X + \lambda V)$. Show that the corresponding Markov kernel can be written as a mixture of Gibbs sampling kernels, and deduce that under the assumptions of Theorem 10.2.7, the hit-and-run sampler admits an entropy contraction bound with rate $1 - 1/\kappa d$.

▷ **Exercise 10.11** (Random subspace hit-and-run)

Generalize the preceding exercise to the setting where, instead of choosing a random line, we choose a random ℓ -dimensional subspace. Show that the entropy contraction factor now becomes $1 - \ell/\kappa d$.

Hint: Write the corresponding Markov kernel as a mixture of “multi-block” Gibbs sampling kernels which simultaneously update ℓ blocks at each iteration.

Mirror Langevin

▷ **Exercise 10.12** (Mirror Langevin diffusion in the primal space)

Use Itô’s formula (Theorem 1.1.19) to show that the mirror Langevin diffusion (10.3.3) in the primal space solves the SDE

$$\begin{aligned} dZ_t = \{ & -[\nabla^2 \phi(Z_t)]^{-1} \nabla V(Z_t) - [\nabla^2 \phi(Z_t)]^{-1} \nabla^3 \phi(Z_t) [\nabla^2 \phi(Z_t)]^{-1} \} dt \\ & + \sqrt{2} [\nabla^2 \phi(Z_t)]^{-1/2} dB_t. \end{aligned}$$

▷ **Exercise 10.13** (Markov semigroup theory for the mirror Langevin diffusion)

Here, we introduce the Markov semigroup perspective on the mirror Langevin diffusion.

- 1 Compute the generator of the mirror Langevin diffusion. Use this to show that π is stationary for the diffusion, and verify the equations (10.3.4) for the carré du champ and Dirichlet energy.
- 2 Let $\mathcal{L}_{\text{dual}}$ denote the generator for $(Z_t^*)_{t \geq 0}$ (we write $\mathcal{L}_{\text{dual}}$ instead of $\mathcal{L}^*$ to avoid confusion with the adjoint of $\mathcal{L}$). By computing $\partial_t \mathbb{E} g(Z_t^*)$, show that

$$\mathcal{L}_{\text{dual}} g = \mathcal{L}(g \circ \nabla \phi) \circ (\nabla \phi)^{-1}.$$

Then, via a similar calculation to (10.3.5) and (10.3.6), show that the Dirichlet energy for $(Z_t^*)_{t \geq 0}$ can be expressed as

$$\mathcal{E}_{\text{dual}}(f, g) = \int \langle \nabla f, [\nabla^2 \phi^*]^{-1} \nabla g \rangle d\pi^*,$$

where $\pi^* := (\nabla \phi)_\# \pi$ is the stationary distribution of $(Z_t^*)_{t \geq 0}$.

- 3 Show that the mirror Poincaré inequality in Corollary 10.3.10 implies the following Poincaré

inequality in the dual space: for all $f : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\text{var}_{\pi^*} f \leq \frac{1}{\alpha} \mathbb{E}_{\pi^*} \langle \nabla f, [\nabla^2 \phi^*]^{-1} \nabla f \rangle = \frac{1}{\alpha} \mathcal{E}_{\text{dual}}(f, f).$$

▷ **Exercise 10.14** (Newton–Langevin diffusion)

Verify the SDE (10.3.7) for the Newton–Langevin diffusion. What happens to the mirror descent iteration (10.3.2) when $\phi = V$?

▷ **Exercise 10.15** (Affine invariance of Newton’s method)

Verify the affine invariance of Newton’s algorithm.

▷ **Exercise 10.16** (Properties of the Bregman divergence)

In this exercise, we check basic properties of the Bregman divergence.

- 1 Prove the alternative definition of relative convexity/smoothness ([Lemma 10.3.9](#)).
- 2 If ϕ^* is the convex conjugate of ϕ , prove that $D_\phi(x, x') = D_{\phi^*}(\nabla \phi(x'), \nabla \phi(x))$.
- 3 Check the identity (10.3.12).

▷ **Exercise 10.17** (Generalized geodesic convexity of the entropy)

Generalize the proof of (1.4.4) to prove the convexity of entropy along interpolations of the Knothe–Rosenblatt map (see the proof of [Theorem 10.2.7](#)), as well as along generalized geodesics ([Theorem 10.3.18](#)).

▷ **Exercise 10.18** (Non-smooth guarantee for MLMC)

Adapt the proof of [Theorem 4.3.11](#) using the techniques of this chapter to prove the non-smooth guarantee for MLMC ([Theorem 10.3.28](#)).

In this chapter, we study the problem of sampling from a smooth but non-log-concave target. Although some results from previous chapters also cover some non-log-concave targets (such as targets satisfying a Poincaré or log-Sobolev inequality), these results do not encompass the full breadth of the non-log-concave sampling problem.

In general, one cannot hope for polynomial-time guarantees from sampling from non-log-concave targets in usual metrics such as total variation distance. Instead, taking inspiration from the literature on non-convex optimization, we will develop a notion of approximate first-order stationarity for sampling, and show that this goal can be achieved via an averaged version of the [LMC](#) algorithm. This is based on the work [Balasubramanian et al. \(2022\)](#).

11.1 Approximate First-Order Stationarity via Fisher Information

Suppose that V is smooth, but non-convex. In general, optimization lower bounds show that finding an approximate global minimizer of V is computationally intractable, i.e., the oracle complexity scales exponentially in the dimension d . To circumvent this, the notion of *approximate first-order stationarity* has arisen as the performance metric of choice in the non-convex optimization literature. Under this metric, we seek the minimal number of queries required to output a point x such that $\|\nabla V(x)\| \leq \varepsilon$.

Optimization Box 11.1.1. Suppose that V is β -smooth and let $\Delta_0 := V(x_0) - \inf V < \infty$. Then, the descent lemma (4.2.2) shows that GD with step size $1/\beta$ satisfies

$$V(x_{n+1}) - V(x_n) \leq -\frac{1}{2\beta} \|\nabla V(x_n)\|^2.$$

Iterating this shows that $\min_{n=0,1,\dots,N-1} \|\nabla V(x_n)\|^2 \leq 2\beta\Delta_0/N$. In particular, we can find an ε -stationary point in $O(\beta\Delta_0/\varepsilon^2)$ iterations, and this rate is sharp in high dimension ([Carmon et al., 2020](#)).

Of course, in practice we may desire stronger guarantees, but first-order stationarity is often a

useful first step towards more detailed analysis, and it has the advantage that we can develop a general theory surrounding this notion. Note that in the convex case, finding a global minimizer is equivalent to finding a first-order stationarity point, so stationary point analysis can be viewed as a natural generalization of the convex optimization analysis to non-convex settings.

To develop a sampling analogue of this concept, we recall that the Langevin diffusion is the gradient flow of the KL divergence $\text{KL}(\cdot \parallel \pi)$ w.r.t. the Wasserstein geometry (Section 1.4). Moreover, the gradient of the KL divergence at μ is $\nabla \log(\mu/\pi)$, and the squared norm of the gradient is the Fisher information $\text{FI}(\mu \parallel \pi) = \mathbb{E}_\mu[\|\nabla \log(\mu/\pi)\|^2]$. Hence, a reasonable definition of finding an approximate first-order stationary point in sampling is to output a sample from μ satisfying $\sqrt{\text{FI}(\mu \parallel \pi)} \leq \varepsilon$. We will show shortly that it is indeed possible to achieve this goal in polynomially many queries to ∇V as soon as ∇V is Lipschitz, thereby establishing a framework for stationarity analysis in non-log-concave sampling. Before doing so, however, we pause to gain intuition for this solution concept.

Lack of spurious stationary points.

An interesting feature of the Fisher information is that, unlike for general non-convex optimization, if π satisfies some mild regularity conditions (e.g., π has a smooth and positive density on $\mathbb{R}^d$), then there are no spurious stationary points: $\text{FI}(\mu \parallel \pi) = 0$ implies $\mu = \pi$. This is a specific feature of the sampling problem.¹ The intuition behind the proof is straightforward: if $\text{FI}(\mu \parallel \pi) = 0$, then $\nabla \log(\mu/\pi) = 0$ (π -a.e.), so the density μ is proportional to π . Since μ is a probability measure, then μ must equal π . (See however the technical remark below.)

This might suggest that our goal of obtaining $\sqrt{\text{FI}(\mu \parallel \pi)} \leq \varepsilon$ is too ambitious, as this would solve the general problem of non-log-concave sampling. This is in fact not the case, and the devil is in the details: while it is true that a small value of $\text{FI}(\mu \parallel \pi)$ should typically imply that μ is close to π , *how small* must $\text{FI}(\mu \parallel \pi)$ be? For highly non-log-concave targets π , typically the Fisher information should be *exponentially small* in order for μ to be close to π in total variation distance.

We illustrate this point with an example: suppose that the target distribution π is a mixture of Gaussians in one dimension, $\pi = \frac{1}{2} \pi_- + \frac{1}{2} \pi_+$, where $\pi_\mp := \text{normal}(\mp m, 1)$. Also, suppose that μ is a mixture of the same two Gaussians, but with the wrong mixing weights: $\mu = \frac{3}{4} \pi_- + \frac{1}{4} \pi_+$. Then, we leave the following computation to the reader (Exercise 11.1).

Proposition 11.1.2. *Let $m > 0$ and let $\pi_\mp := \text{normal}(\mp m, 1)$. Let $\mu := \frac{3}{4} \pi_- + \frac{1}{4} \pi_+$ and $\pi := \frac{1}{2} \pi_- + \frac{1}{2} \pi_+$. Then, it holds that*

$$\liminf_{m \rightarrow \infty} \|\mu - \pi\|_{\text{TV}} > 0$$

whereas

$$\text{FI}(\mu \parallel \pi) \lesssim m^2 \exp\left(-\frac{m^2}{2}\right) \rightarrow 0 \quad \text{as } m \rightarrow \infty.$$

Metastability.

The example in Proposition 11.1.2 also provides an interpretation of the Fisher information. If we try to sample from the mixture of Gaussians π , then for $m \gg 1$ it takes an exponentially long time for the Langevin diffusion to jump to one mode from the other; this is the main reason behind the slow mixing

¹ In fact, the KL divergence $\text{KL}(\cdot \parallel \pi)$ is always strictly convex with respect to taking convex combinations of measures, and hence always has a unique global minimum.

of Langevin. Since it is hard to jump between the modes, it is difficult for the Langevin diffusion to “learn” the global mixing weights $(\frac{1}{2}, \frac{1}{2})$ of π . On the other hand, [Proposition 11.1.2](#) shows that even with the wrong mixing weights $(\frac{3}{4}, \frac{1}{4})$, the Fisher information $\text{FI}(\mu \parallel \pi)$ is small, demonstrating that the Fisher information is insensitive to the global weights.

The example in [Proposition 11.1.2](#) therefore paints a cartoon picture of the behavior of the Langevin diffusion with target π , initialized at $\frac{3}{4} \delta_{-m} + \frac{1}{4} \delta_{+m}$: we expect that the Langevin diffusion quickly explores and captures the local structure of the modes but fails to jump between the modes, arriving at a distribution which resembles μ ; it is this local mixing that a Fisher information bound captures. Meanwhile, the Langevin diffusion only obtains the correct global weights after an exponentially long waiting time. This is further formalized in [Theorem 11.3.8](#) below.

The state μ is not truly stable for the Langevin diffusion: given enough time, the diffusion will eventually move away from μ and reach π . However, since states like μ persist for a very long period of time, they are usually called *metastable* in the statistical physics literature. A Fisher information bound can be interpreted as a way of quantitatively measuring the metastability phenomenon.

Technical remark.

One has to be slightly careful with the definition of the Fisher information. For example, suppose that π is the standard Gaussian, and suppose that μ is the Gaussian restricted to the unit ball. Then, it is tempting to argue that the density of μ is proportional to that of π on the unit ball, and hence $\nabla \log(\mu/\pi) = 0$ (μ -a.e.); from the expression $\text{FI}(\mu \parallel \pi) = \mathbb{E}_\mu[\|\nabla \log(\mu/\pi)\|^2]$, it suggests that $\text{FI}(\mu \parallel \pi) = 0$, and in particular μ is a spurious stationary point. However, this argument is *not* correct.

The reason is that μ does not have enough regularity w.r.t. π in order to apply the formula $\text{FI}(\mu \parallel \pi) = \mathbb{E}_\mu[\|\nabla \log(\mu/\pi)\|^2]$. Indeed, in order to apply the formula, we must require that the density $\frac{d\mu}{d\pi}$ lie in an appropriate Sobolev space w.r.t. π (more precisely, $(\frac{d\mu}{d\pi})^{1/2}$ should lie in the domain of the Dirichlet energy functional). If this does not hold, then we define the Fisher information to be infinite: $\text{FI}(\mu \parallel \pi) = \infty$. In our results below, the Fisher information bound should be interpreted as follows: $\sqrt{\text{FI}(\mu \parallel \pi)} \leq \varepsilon$ means that μ has enough regularity w.r.t. π and $\mathbb{E}_\mu[\|\nabla \log(\mu/\pi)\|^2] \leq \varepsilon^2$.

11.2 Fisher Information Bounds

As in [Section 4.2](#), we consider the interpolation of [LMC](#),

$$\hat{X}_t := \hat{X}_{nh} - (t - nh) \nabla V(\hat{X}_{nh}) + \sqrt{2} (B_t - B_{nh}), \quad t \in [nh, (n+1)h].$$

Theorem 11.2.1 ([Balasubramanian et al. \(2022\)](#)). *Let $(\hat{\mu}_t)_{t \geq 0}$ denote the law of the interpolation of [LMC](#) with step size $h > 0$. Assume that $\pi \propto \exp(-V)$ where ∇V is β -Lipschitz. Then, for any step size $0 < h \leq 1/4\beta$, for all $N \in \mathbb{N}$,*

$$\frac{1}{Nh} \int_0^{Nh} \text{FI}(\hat{\mu}_t \parallel \pi) dt \leq \frac{2 \text{KL}(\hat{\mu}_0 \parallel \pi)}{Nh} + 6\beta^2 dh.$$

In particular, if $\text{KL}(\hat{\mu}_0 \parallel \pi) \leq K_0$ and we choose $h = \sqrt{K_0}/(2\beta\sqrt{dN})$, then provided that $N \geq 9K_0/d$,

$$\frac{1}{Nh} \int_0^{Nh} \text{FI}(\hat{\mu}_t \parallel \pi) dt \leq \frac{8\beta\sqrt{dK_0}}{\sqrt{N}}.$$

In order to translate the result into a more useful form, we recall that the Fisher information is convex in its first argument.

Lemma 11.2.2. *The Fisher information functional $\text{FI}(\cdot \parallel \pi)$ is convex.*

Proof Let $\mu_0, \mu_1 \in \mathcal{P}(\mathbb{R}^d)$ be such that $\text{FI}(\mu_0 \parallel \pi) \vee \text{FI}(\mu_1 \parallel \pi) < \infty$. For $t \in (0, 1)$, let $\mu_t := (1 - t)\mu_0 + t\mu_1$, and write $f_t := \frac{d\mu_t}{d\pi} = (1 - t)f_0 + tf_1$. Then,

$$\begin{aligned} \text{FI}(\mu_t \parallel \pi) &= \int \frac{\|\nabla f_t\|^2}{f_t} d\pi \leq (1 - t) \int \frac{\|\nabla f_0\|^2}{f_0} d\pi + t \int \frac{\|\nabla f_1\|^2}{f_1} d\pi \\ &= (1 - t) \text{FI}(\mu_0 \parallel \pi) + t \text{FI}(\mu_1 \parallel \pi) \end{aligned}$$

follows from the joint convexity of $(a, b) \mapsto \|a\|^2/b$ on $\mathbb{R}^d \times \mathbb{R}_{>0}$. $\square$

Hence, for the averaged measure $\bar{\mu}_{Nh} := \frac{1}{Nh} \int_0^{Nh} \hat{\mu}_t dt$, we have

$$\text{FI}(\bar{\mu}_{Nh} \parallel \pi) \leq \frac{1}{Nh} \int \text{FI}(\hat{\mu}_t \parallel \pi) dt \quad (11.2.3)$$

and the guarantees of the theorem translate into guarantees for $\bar{\mu}_{Nh}$. Moreover, we can output a sample from $\bar{\mu}_{Nh}$ via the following procedure:

- 1 Pick a time $t \in [0, Nh]$ uniformly at random.
- 2 Let n be the largest integer such that $nh \leq t$, and let $\hat{X}_{nh}$ denote the n -th iterate of **LMC**.
- 3 Perform a partial **LMC** update

$$\hat{X}_t = \hat{X}_{nh} - (t - nh) \nabla V(\hat{X}_{nh}) + \sqrt{2} (B_t - B_{nh})$$

and output $\hat{X}_t$.

Combined with [Theorem 11.2.1](#) and (11.2.3), we conclude that it is possible to algorithmically obtain a sample from a measure $\hat{\mu}$ with $\sqrt{\text{FI}(\hat{\mu} \parallel \pi)} \leq \varepsilon$ using $O(\beta^2 d K_0 / \varepsilon^4)$ queries to ∇V .

We now give the proof of [Theorem 11.2.1](#), which combines the usual stationary point analysis in non-convex optimization with the interpolation argument.

Proof of Theorem 11.2.1 Recall from the proof of [Theorem 4.2.7](#) that

$$\partial_t \text{KL}(\hat{\mu}_t \parallel \pi) \leq -\frac{1}{2} \text{FI}(\hat{\mu}_t \parallel \pi) + 6\beta^2 d (t - nh).$$

This inequality was obtained under the sole assumption that ∇V is β -Lipschitz. In [Theorem 4.2.7](#), we proceeded to apply a log-Sobolev inequality, but here we will instead telescope this inequality. By integrating over $t \in [nh, (n+1)h]$,

$$\text{KL}(\hat{\mu}_{(n+1)h} \parallel \pi) - \text{KL}(\hat{\mu}_{nh} \parallel \pi) \leq -\frac{1}{2} \int_{nh}^{(n+1)h} \text{FI}(\hat{\mu}_t \parallel \pi) dt + 3\beta^2 dh^2.$$

Summing over $n = 0, 1, \dots, N-1$ and dividing by Nh ,

$$\frac{1}{Nh} \int_0^{Nh} \text{FI}(\hat{\mu}_t \parallel \pi) dt \leq \frac{2 \text{KL}(\hat{\mu}_0 \parallel \pi)}{Nh} + 6\beta^2 dh.$$

This proves the first statement; the second statement is obtained by optimizing over h . $\square$

Although [Theorem 11.2.1](#) shares some similarities with the corresponding result from optimization ([Optimization Box 11.1.1](#)), the rate leaves room for improvement. It can be shown that the proximal sampler (Chapter 8) also inherits Fisher information decay from the Langevin diffusion ([Exercise 11.3](#)). Together with some careful bookkeeping regarding RGO implementation, [Chewi and Wibisono \(2026\)](#) established the following reduction to high-accuracy log-concave sampling.

Theorem 11.2.4 ([Chewi and Wibisono \(2026\)](#)). *Suppose that there is a sampler with the following guarantee: given query access to a strongly log-concave and smooth distribution with mode at the origin and condition number at most 3, the sampler returns a δ -accurate sample in $\mathcal{R}_3$ using $\tilde{O}(d^p \text{polylog}(1/\delta))$ queries.*

Then, there is a sampler which yields a sample from a distribution $\hat{\pi}$ with $\text{FI}(\hat{\pi} \parallel \pi) \leq \varepsilon^2$ using

$$\tilde{O}\left(\frac{\beta d^p K_0}{\varepsilon^2} \text{polylog}(R_0, d, \beta/\varepsilon^2)\right) \quad \text{queries,}$$

where $K_0 := \text{KL}(\mu_0 \parallel \pi)$ and $R_0 := \mathcal{R}_2(\mu_0 \parallel \pi)$.

In particular, by extending either [Corollary 7.3.6](#) or [Theorem 8.6.6](#) to hold in $\mathcal{R}_3$, the assumption is met with $p = 1/2$, yielding a complexity of $\tilde{O}(\beta d^{1/2} K_0 / \varepsilon^2)$. Unfortunately, the result also depends mildly on the additional quantity R_0 , but this seems to be a technical defect rather than a fundamental one. Interestingly, the dimension dependence here cannot be improved without also improving the dimension dependence of high-accuracy log-concave sampling; see [Exercise 11.4](#).

11.3 Applications of Fisher Information Bounds

Asymptotic convergence of averaged LMC.

Since [Theorem 11.2.1](#) holds under very weak assumptions (only smoothness of V) and implies that the Fisher information can be driven to zero, and $\text{FI}(\mu \parallel \pi) = 0$ implies that $\mu = \pi$, putting these facts together leads to a straightforward proof of the asymptotic convergence of averaged LMC.

For a sequence of positive step sizes $(h_n)_{n \in \mathbb{N}^+}$, let $\tau_n := \sum_{k=1}^n h_k$ and consider the interpolation

$$\hat{X}_t = \hat{X}_{\tau_{n-1}} - (t - \tau_{n-1}) \nabla V(\hat{X}_{\tau_{n-1}}) + \sqrt{2} (B_t - B_{\tau_{n-1}}), \quad t \in [\tau_{n-1}, \tau_n]. \quad (11.3.1)$$

Theorem 11.3.2. *Let $(\hat{\mu}_t)_{t \geq 0}$ denote the law of the interpolation (11.3.1) of LMC, and suppose that the target is $\pi \propto \exp(-V)$ where ∇V is β -Lipschitz. Suppose that we initialize at a measure $\hat{\mu}_0$ with $\text{KL}(\hat{\mu}_0 \parallel \pi) < \infty$, and that the sequence of step sizes satisfies $0 < h_n < 1/4\beta$ for all $n \in \mathbb{N}^+$ together with the conditions*

$$\sum_{n=1}^{\infty} h_n = \infty, \quad \text{and} \quad \sum_{n=1}^{\infty} h_n^2 < \infty.$$

Write $\bar{\mu}_{\tau_n} := \frac{1}{\tau_n} \int_0^{\tau_n} \hat{\mu}_t dt$. Then, $\bar{\mu}_{\tau_n} \rightarrow \pi$ weakly.

Proof By repeating the proof of [Theorem 11.2.1](#) but incorporating time-varying step sizes, we obtain

for $t \in [\tau_n, \tau_{n+1}]$

$$\partial_t \text{KL}(\hat{\mu}_t \parallel \pi) \leq -\frac{1}{2} \text{FI}(\hat{\mu}_t \parallel \pi) + 6\beta^2 d (t - \tau_n). \quad (11.3.3)$$

By integrating this inequality and summing,

$$\text{KL}(\hat{\mu}_{\tau_n} \parallel \pi) \leq \text{KL}(\hat{\mu}_0 \parallel \pi) - \frac{1}{2} \int_0^{\tau_n} \text{FI}(\hat{\mu}_t \parallel \pi) dt + 3\beta^2 d \sum_{k=1}^n h_k^2. \quad (11.3.4)$$

Rearranging and using the convexity of the Fisher information, it yields

$$\text{FI}(\bar{\mu}_{\tau_n} \parallel \pi) \leq \frac{1}{\tau_n} \int_0^{\tau_n} \text{FI}(\hat{\mu}_t \parallel \pi) dt \leq \frac{2 \text{KL}(\hat{\mu}_0 \parallel \pi)}{\tau_n} + \frac{6\beta^2 d}{\tau_n} \sum_{k=1}^n h_k^2. \quad (11.3.5)$$

On the other hand, if $t \in [\tau_n, \tau_{n+1}]$, then integrating (11.3.3) and combining with (11.3.4) yields

$$\text{KL}(\hat{\mu}_t \parallel \pi) \leq \text{KL}(\hat{\mu}_{\tau_n} \parallel \pi) + 3\beta^2 d (t - \tau_n)^2 \leq \text{KL}(\hat{\mu}_0 \parallel \pi) + 6\beta^2 d \sum_{k=1}^{\infty} h_k^2 < \infty.$$

Therefore, $\{\text{KL}(\hat{\mu}_t \parallel \pi) \mid t \geq 0\}$ is bounded, and the convexity of the KL divergence implies that $\{\text{KL}(\bar{\mu}_{\tau_n} \parallel \pi) \mid n \in \mathbb{N}^+\}$ is bounded. Since the sublevel sets of $\text{KL}(\cdot \parallel \pi)$ are compact, to prove the theorem it suffices to show that every weak limit of $(\bar{\mu}_{\tau_n})_{n \in \mathbb{N}^+}$ is equal to π . Consider a subsequence of $(\bar{\mu}_{\tau_n})_{n \in \mathbb{N}^+}$ converging to a weak limit $\bar{\mu}$.

Taking $n \rightarrow \infty$ in (11.3.5) and noting that $\tau_n \rightarrow \infty$, we have $\text{FI}(\bar{\mu}_{\tau_n} \parallel \pi) \rightarrow 0$ and thus along the subsequence as well. It is known that $\text{FI}(\cdot \parallel \pi)$ is weakly lower semicontinuous, so $\text{FI}(\bar{\mu} \parallel \pi) = 0$. However, since ∇V is Lipschitz, then π has a continuous and strictly positive density on $\mathbb{R}^d$, so $\text{FI}(\bar{\mu} \parallel \pi) = 0$ entails $\bar{\mu} = \pi$ as desired. $\square$

Convergence in total variation distance under a Poincaré inequality.

As the example in Proposition 11.1.2 shows, for general non-log-concave targets, a Fisher information bound does not translate into guarantees in other metrics. However, this can be carried out if π satisfies appropriate functional inequalities. For example, by definition, a log-Sobolev inequality for π yields

$$\text{KL}(\mu \parallel \pi) \leq \frac{C_{\text{LSI}}(\pi)}{2} \text{FI}(\mu \parallel \pi) \quad \text{for all } \mu \in \mathcal{P}(\mathbb{R}^d),$$

and in this case a Fisher information guarantee readily translates into a KL divergence guarantee; however, this is not very interesting because we have obtained a sharper KL divergence guarantee for targets π satisfying a log-Sobolev inequality in Theorem 4.2.7. Instead, we will show that under the weaker assumption of a Poincaré inequality, a Fisher information guarantee implies a TV guarantee.

The key observation is the following implication of a Poincaré inequality.

Proposition 11.3.6. *Suppose that π satisfies a Poincaré inequality with constant C_{PI} . Then, for all $\mu \in \mathcal{P}(\mathbb{R}^d)$,*

$$\|\mu - \pi\|_{\text{TV}}^2 \leq \frac{C_{\text{PI}}}{4} \text{FI}(\mu \parallel \pi).$$

Proof We can assume $\mu \ll \pi$; let $f := \frac{d\mu}{d\pi}$. The total variation distance has the expressions

$$\|\mu - \pi\|_{\text{TV}} = \frac{1}{2} \int |f - 1| d\pi = \frac{1}{2} \int \{(f \vee 1) - (f \wedge 1)\} d\pi$$

which yields $\int (f \wedge 1) d\pi = 1 - \|\mu - \pi\|_{\text{TV}}$ and $\int (f \vee 1) d\pi = 1 + \|\mu - \pi\|_{\text{TV}}$. Using this,

$$\begin{aligned} \int \sqrt{f} d\pi &= \int \sqrt{(f \wedge 1)(f \vee 1)} d\pi \leq \sqrt{\int (f \wedge 1) d\pi} \sqrt{\int (f \vee 1) d\pi} \\ &= \sqrt{(1 - \|\mu - \pi\|_{\text{TV}})(1 + \|\mu - \pi\|_{\text{TV}})} = \sqrt{1 - \|\mu - \pi\|_{\text{TV}}^2}. \end{aligned}$$

Therefore,

$$\|\mu - \pi\|_{\text{TV}}^2 \leq 1 - \left(\int \sqrt{f} d\pi \right)^2 = \text{var}_\pi \sqrt{f}.$$

This is sometimes called **Le Cam's inequality**; in statistics, the right-hand side is often written as $H^2(\mu, \pi) (1 - \frac{1}{4} H^2(\mu, \pi))$, where H^2 denotes the squared **Hellinger distance**.

Applying the Poincaré inequality,

$$\|\mu - \pi\|_{\text{TV}}^2 \leq C_{\text{PI}} \mathbb{E}_\pi [\|\nabla \sqrt{f}\|^2] = \frac{C_{\text{PI}}}{4} \text{FI}(\mu \parallel \pi). \quad \square$$

Corollary 11.3.7. Let $(\hat{\mu}_t)_{t \geq 0}$ denote the law of the interpolation of **LMC** with step size $h > 0$. Assume that $\pi \propto \exp(-V)$, where ∇V is β -Lipschitz and π satisfies the Poincaré inequality with constant C_{PI} . If $\text{KL}(\hat{\mu}_0 \parallel \pi) \leq K_0$ and we choose step size $h = \sqrt{K_0}/(2\beta\sqrt{dN})$, then

$$\|\bar{\mu}_{Nh} - \pi\|_{\text{TV}}^2 := \left\| \frac{1}{Nh} \int_0^t \hat{\mu}_t dt - \pi \right\|_{\text{TV}}^2 \leq \frac{2C_{\text{PI}}\beta\sqrt{dK_0}}{\sqrt{N}}.$$

Reweighted functional inequalities.

Recall the intuition from the example of [Proposition 11.1.2](#): although Langevin mixes rapidly within each component, the barrier to global equilibration lies in the mixture weights. This idea was elegantly formulated via **reweighted functional inequalities** in [Niles-Weed \(2026\)](#).

Theorem 11.3.8 (Reweight functional inequalities, [Niles-Weed \(2026\)](#)). Let $\mu, \pi_1, \dots, \pi_K$ be probability measures. Suppose that $\pi \in \mathcal{C} := \text{conv}\{\pi_1, \dots, \pi_K\}$.

1 If each π_k satisfies a log-Sobolev inequality with constant C_{LSI} , then

$$\inf_{\tilde{\pi} \in \mathcal{C}} \text{KL}(\mu \parallel \tilde{\pi}) \leq \frac{C_{\text{LSI}}}{2} \text{FI}(\mu \parallel \pi).$$

2 If each π_k satisfies a Poincaré inequality with constant C_{PI} , then

$$\inf_{\tilde{\pi} \in \mathcal{C}} \|\mu - \tilde{\pi}\|_{\text{TV}}^2 \leq \frac{C_{\text{PI}}}{4} \text{FI}(\mu \parallel \pi).$$

The interpretation of these results is that if π is a mixture of $\pi_1, \dots, \pi_K$, where each component π_k satisfies a functional inequality, then the Fisher information guarantee $\text{FI}(\mu \parallel \pi) \leq \varepsilon^2$ implies that μ

is close, in the sense of KL or TV, to some other $\tilde{\pi} \in \mathcal{C}$. Since $\tilde{\pi}$ is a mixture of the same components $\pi_1, \dots, \pi_K$, but with potentially different mixing weights, we refer to $\tilde{\pi}$ as a *reweighting* of π .

Proof Write $\pi := \sum_{k=1}^K p_k \pi_k$ and let $\tilde{p}_k := p_k \int \frac{d\mu}{d\pi} d\pi_k$. The key fact is that if we set $\frac{d\mu_k}{d\pi_k} := \frac{p_k}{\tilde{p}_k} \frac{d\mu}{d\pi}$, then $\sum_{k=1}^K \tilde{p}_k \mu_k = \mu$, and in particular, $\sum_{k=1}^K \tilde{p}_k = 1$. We will show that the bounds hold for the particular element $\tilde{\pi} := \sum_{k=1}^K \tilde{p}_k \pi_k \in \mathcal{C}$.

It follows from joint convexity that

$$\begin{aligned} \text{KL}(\mu \parallel \tilde{\pi}) &\leq \sum_{k=1}^K \tilde{p}_k \text{KL}(\mu_k \parallel \pi_k) \leq \frac{C_{\text{LSI}}}{2} \sum_{k=1}^K \tilde{p}_k \text{FI}(\mu_k \parallel \pi_k) = \frac{C_{\text{LSI}}}{2} \sum_{k=1}^K \tilde{p}_k \int \left\| \nabla \log \frac{d\mu_k}{d\pi_k} \right\|^2 d\mu_k \\ &= \frac{C_{\text{LSI}}}{2} \text{FI}(\mu \parallel \pi). \end{aligned}$$

We leave the proof of the second inequality as [Exercise 11.5](#). □

11.4 Lower Bounds

We now turn toward lower bounds, following [Chewi et al. \(2023a\)](#). Let $\mathcal{C}(d, K_0, \varepsilon; \beta)$ denote the complexity of achieving $\sqrt{\text{FI}} \leq \varepsilon$ in dimension d for a β -log-smooth target, and with initial KL divergence at most K_0 . The first observation is that the complexity only depends on the ratio $\varepsilon/\sqrt{\beta}$.

Lemma 11.4.1. *It holds that $\mathcal{C}(d, K_0, \varepsilon; \beta) = \mathcal{C}(d, K_0, \varepsilon/\sqrt{\beta}; 1)$.*

See [Exercise 11.6](#). In particular, we can set $\varepsilon = 1$ for ease of notation and write $\mathcal{C}(d, K_0, \varepsilon) := \mathcal{C}(d, K_0, \varepsilon; 1)$. The lower bounds for the general case are obtained by replacing ε by $\varepsilon/\sqrt{\beta}$ below.

The first lower bound is proven for the regime $\varepsilon \asymp \sqrt{d}$ via a reduction to finding stationary points in non-convex optimization. The key is the following lemma, which is left as [Exercise 11.7](#).

Lemma 11.4.2. *Let $V : \mathbb{R}^d \rightarrow \mathbb{R}$ be 1-smooth, such that for any $\beta > 0$, $\int \exp(-\beta V) < \infty$. Let $\pi_\beta \propto \exp(-\beta V)$, where $\beta = d/\varepsilon^2$.*

- 1 *If $\|\nabla V(x)\| \leq \varepsilon$, then $\text{FI}(\text{normal}(x, \beta^{-1} I_d) \parallel \pi_\beta) \leq 10\beta d$.*
- 2 *Conversely, if $\text{FI}(\mu \parallel \pi_\beta) \leq \beta d$, then for $X \sim \mu$, we have $\|\nabla V(X)\| \leq 3\varepsilon$ with probability at least $1/2$.*

This lemma shows that finding an ε -stationary point for V is equivalent to finding a measure with $\text{FI}(\cdot \parallel \pi_\beta) \leq \beta d$ for $\beta = d/\varepsilon^2$, up to numerical constants. Using this, let us calculate the complexity of finding an ε -stationary point for V using averaged [LMC](#). Letting $\Delta_0 := V(x_0) - \inf V$ and $\mu_0 := \text{normal}(x_0, \beta^{-1} I_d)$, one can check that $K_0 := \text{KL}(\mu_0 \parallel \pi_\beta) \leq \tilde{O}(\beta \Delta_0 + d)$. Since π_β is β -smooth, [Theorem 11.2.1](#) implies that averaged [LMC](#) obtains a measure with $\text{FI}(\cdot \parallel \pi_\beta) \leq \beta d$ in $O(\beta^2 d K_0 / \beta^2 d^2) = \tilde{O}(\beta \Delta_0 / d + 1) = \tilde{O}(\Delta_0 / \varepsilon^2 + 1)$ iterations. The details are given as [Exercise 11.8](#). Note that this matches the iteration complexity of gradient descent from [Optimization Box 11.1.1](#). This is unsurprising, since [LMC](#) with potential βV , $\beta \asymp d/\varepsilon^2$, is already quite close to GD.

However, we can now leverage the fact that GD is *optimal* for finding stationary points in high dimension to turn this into a lower bound. In particular, by passing the lower bound of [Carmon et al. \(2020\)](#) into the reduction of [Lemma 11.4.2](#), [Chewi et al. \(2023a\)](#) established:

Theorem 11.4.3 (First lower bound). Assume that $\tilde{\Omega}(K_0^{2/3}) \leq d \leq \tilde{O}(K_0)$. Then,

$$\mathcal{C}(d, K_0, \varepsilon = \sqrt{d}) \asymp \frac{K_0}{d}.$$

The restriction on d arises because the lower bound construction of [Carmon et al. \(2020\)](#) is embedded in very high dimension, and is not likely to be fundamental. Since the lower bound is attained by [Theorem 11.2.1](#), we have identified one particular regime in which we can pin down the Fisher information complexity up to constants.

As for the regime of small ε , since the distributions under consideration are not restricted to be log-concave, this avoids some of the difficulties of the constructions from Chapter 9. In particular, using a careful bump function construction, [Chewi et al. \(2023a\)](#) established:

Theorem 11.4.4 (Second lower bound). The following lower bounds hold for $\varepsilon \ll 1$.

- 1 $\mathcal{C}(d = 1, K_0 = 1, \varepsilon) \gtrsim 1/\varepsilon \sqrt{\log(1/\varepsilon)}$.
- 2 $\mathcal{C}(d, K_0 = 1, \varepsilon) \gtrsim 1/\varepsilon^{2-o(1)}$ for all $d \gtrsim \sqrt{\log(1/\varepsilon)/\log \log(1/\varepsilon)}$.

Comparing with [Theorem 11.2.4](#), we see that the $1/\varepsilon^2$ dependence on the accuracy parameter ε is optimal in high dimension. However, [Theorem 11.4.4](#) does not capture dependence on the parameters d, K_0 . Intriguingly, due to the reduction of [Theorem 11.2.4](#), if [Theorem 11.4.4](#) could be improved to exhibit a polynomial dependence on d , it would imply the same polynomial dimension dependence for high-accuracy log-concave sampling, which is currently open.

Recall from [Optimization Box 11.1.1](#) that the optimal complexity for finding stationary points in optimization is $\Theta(\beta \Delta_0 / \varepsilon^2)$ in high dimension. In sampling, K_0 plays the role of Δ_0 , the initial gap. In light of [Theorem 11.2.4](#), [Exercise 11.4](#), and [Theorem 11.4.4](#), it seems reasonable to conjecture that the optimal FI complexity in high dimension is $\Theta(\beta d^p K_0 / \varepsilon^2)$, where d^p is the (unknown) optimal dimensional complexity of high-accuracy log-concave sampling.

Bibliographical Notes

The framework for non-log-concave sampling, as well as many of the results from this chapter, are from [Balasubramanian et al. \(2022\)](#). The guarantee of [Theorem 11.2.1](#) was improved in [Zhou and Sugiyama \(2025a\)](#) via parallelization, and extended to mirror Langevin in [Daaloul et al. \(2025\)](#) and to the zeroth-order setting in [Sahin et al. \(2026\)](#). The lower bounds are from [Chewi et al. \(2023a\)](#).

In a similar spirit to [Theorem 11.3.8](#), [Cheng et al. \(2023\)](#) proved that for mixture of Gaussian targets, $\text{FI}(\mu \parallel \pi) \leq \varepsilon^2$ implies mixing for each of the components which have non-trivial mass under μ .

The formula in [Exercise 11.2](#) is from [Wibisono \(2025\)](#), and [Exercise 11.3](#) and [Exercise 11.4](#) are from [Chewi and Wibisono \(2026\)](#).

Exercises

Approximate First-Order Stationarity via Fisher Information

▷ **Exercise 11.1** (Fisher information for mixtures of Gaussians)

Prove Proposition 11.1.2.

Fisher Information Bounds

▷ **Exercise 11.2** (Fisher information decay along simultaneous flows)

Suppose that $(\mu_t)_{t \geq 0}$, $(\nu_t)_{t \geq 0}$ both evolve according to the Fokker–Planck equations

$$\partial_t \mu_t = \operatorname{div}(\mu_t \nabla V_t) + \frac{c}{2} \Delta \mu_t, \quad \partial_t \nu_t = \operatorname{div}(\nu_t \nabla V_t) + \frac{c}{2} \Delta \nu_t.$$

Establish the formula

$$\partial_t \operatorname{FI}(\mu_t \parallel \nu_t) = -c \mathbb{E}_{\mu_t} \left[\left\| \nabla^2 \log \frac{\mu_t}{\nu_t} \right\|_{\text{HS}}^2 \right] - 2 \mathbb{E}_{\mu_t} \left\langle \nabla \log \frac{\mu_t}{\nu_t}, \left(c \nabla^2 \log \frac{1}{\nu_t} - \nabla^2 V_t \right) \nabla \log \frac{\mu_t}{\nu_t} \right\rangle.$$

Hint: Either proceed via a direct (but tedious) time derivative calculation, or follow the approach of [Exercise 3.5](#).

▷ **Exercise 11.3** (Fisher information decay along the proximal sampler)

In the notation of [Chapter 8](#), prove the following inequalities for $h \leq 1/\beta$.

- 1 $h(1 - \beta h) \operatorname{FI}(\mu_0^Y \parallel \pi^Y) \leq 2 \{ \operatorname{KL}(\mu_0^X \parallel \pi^X) - \operatorname{KL}(\mu_0^Y \parallel \pi^Y) \}.$
- 2 $h \operatorname{FI}(\mu_1^X \parallel \pi^X) \leq 2 \{ \operatorname{KL}(\mu_0^Y \parallel \pi^Y) - \operatorname{KL}(\mu_1^X \parallel \pi^X) \}.$

Hint: Use [Exercise 2.11](#), [\(8.3.2\)](#), and [Exercise 11.2](#).

▷ **Exercise 11.4** (Converse reduction)

Suppose that there is a sampler **Alg** with the following property: for any β -log-smooth distribution π and initialization μ_0 satisfying $\operatorname{KL}(\mu_0 \parallel \pi) \leq K_0$ and $\mathcal{R}_2(\mu_0 \parallel \pi) \leq R_0$, **Alg** returns a sample from $\hat{\pi}$ with $\mathcal{R}_2(\hat{\pi} \parallel \pi) \leq R_0$ and $\operatorname{FI}(\hat{\pi} \parallel \pi) \leq \varepsilon^2$ using

$$\frac{\beta \phi(d) K_0}{\varepsilon^2} \operatorname{polylog}(R_0, \beta/\varepsilon^2) \quad \text{queries}.$$

Then, this implies the existence of a sampler which, given query access to a strongly log-concave and log-smooth distribution with mode at the origin, outputs a sample from $\hat{\pi}$ satisfying $\operatorname{KL}(\hat{\pi} \parallel \pi) \leq \varepsilon^2$ using $O(\kappa \phi(d) \operatorname{polylog}(d, \kappa, 1/\varepsilon))$ queries. Hence, conclude that [Theorem 11.2.4](#), which reduces the dimension dependence of FI guarantees for non-log-concave sampling to the dimension dependence of high-accuracy log-concave sampling, is optimal.

Hint: Using the log-Sobolev inequality, consider how many iterations of **Alg** are required to reduce the KL divergence to π from K_0 to $K_0/2$.

Applications of Fisher Information Bounds

▷ **Exercise 11.5** (Reweighted functional inequalities)

In the setting of [Theorem 11.3.8](#), let $\mathcal{T}_c$ be a transport cost (see [Definition 1.3.1](#)) and let $\phi : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ be convex, lower semicontinuous, and increasing with $\phi(0) = 0$. Suppose that each π_k satisfies a transport-information inequality of the form

$$\phi(\mathcal{T}_c(\mu, \pi_k)) \leq C \operatorname{FI}(\mu \parallel \pi_k) \quad \text{for all } \mu \in \mathcal{P}(\mathbb{R}^d).$$

Show that

$$\phi\left(\inf_{\tilde{\pi} \in \mathcal{C}} \mathcal{T}_c(\mu, \tilde{\pi})\right) \leq C \operatorname{FI}(\mu \parallel \pi).$$

In particular, if each π_k satisfies a log-Sobolev inequality with constant C_{LSI} , deduce that

$$\inf_{\tilde{\pi} \in \mathcal{C}} W_2^2(\mu, \tilde{\pi}) \leq C_{\text{LSI}}^2 \text{FI}(\mu \parallel \pi).$$

Also, by specializing to $c : (x, y) \mapsto \mathbb{1}_{x \neq y}$, prove the second inequality in [Theorem 11.3.8](#).

Lower Bounds

⊇ **Exercise 11.6** (Rescaling argument)

Prove Lemma 11.4.1.

⊇ **Exercise 11.7** (Reduction to finding stationary points)

Prove Lemma 11.4.2.

⊇ **Exercise 11.8** (Finding stationary points via averaged LMC)

Assume that $\int \|\cdot\|^2 d\pi_\beta \leq \text{poly}(\Delta_0, d, 1/\varepsilon)$. Use [Exercise 6.7](#) to argue that $\text{KL}(\mu_0 \parallel \pi_\beta) = \tilde{O}(\beta\Delta_0 + d)$, where $\mu_0 := \text{normal}(x_0, \beta^{-1}I_d)$. Verify the discussion following [Lemma 11.4.2](#).

In this final chapter, we present a brief introduction to diffusion generative models, focusing on the analysis of their discretization.

12.1 Introduction

Diffusion generative models, since their introduction in [Sohl-Dickstein et al. \(2015\)](#); [Song and Ermon \(2019\)](#); [Ho et al. \(2020\)](#), have quickly emerged as a predominant paradigm for generative modelling. This is a large (and growing) field with numerous important considerations that lie beyond the scope of the book. Our aim here is more modest: namely, to present the interpretation of the basic model as the time reversal of an SDE, and to give a self-contained and relatively refined discretization analysis, building upon the tools developed in previous chapters (most notably, Section 4.4).

We follow the presentation in [Song et al. \(2021b\)](#). As in previous chapters, the goal is to sample from a probability measure π over $\mathbb{R}^d$. Consider running an SDE forward in time starting from π :

$$dX_t^\rightarrow = f_t(X_t^\rightarrow) dt + g_t dB_t, \quad X_0^\rightarrow \sim \pi. \quad (12.1.1)$$

Here, for simplicity, we restrict the diffusion coefficient g_t to be a scalar. Now, fix a terminal time $T > 0$ and for each $t \in [0, T]$, let $\pi_t := \text{law}(X_t^\rightarrow)$. Then, by the theory of Section 3.3.2, the *time reversal* $(X_t^\leftarrow)_{t \in [0, T]} := (X_{T-t}^\rightarrow)_{t \in [0, T]}$ is also described by an SDE,

$$dX_t^\leftarrow = \{-f_{T-t}(X_t^\leftarrow) + g_{T-t}^2 \nabla \log \pi_{T-t}(X_t^\leftarrow)\} dt + g_{T-t} dB_t, \quad X_0^\leftarrow \sim \pi_T. \quad (12.1.2)$$

If we can simulate the reverse SDE (12.1.2) with initialization $X_0^\leftarrow \sim \pi_T$, then this would yield a mechanism to generate a sample from the target distribution: $X_T^\leftarrow \sim \pi$. In other words, a generative model. However, we must first address three implementation details.

- 1 *Initialization.* Typically, we do not have access to the exact initialization π_T , but the forward SDE (12.1.1) can be chosen so that π_T is close to a known distribution γ , which can then serve as the initialization. For example, this is the case if (12.1.1) admits γ as a stationary distribution and T is chosen sufficiently large.

- 2 *Score functions.* Simulation of the reverse SDE (12.1.2) requires evaluation of the *score functions* $\{\nabla \log \pi_t\}_{0 < t \leq T}$ along the diffusion. The perspective that we adopt in this chapter is to view the score functions as the manner in which we have access to the distribution π , and the computational goal is to sample well while minimizing the number of score function evaluations. This is a significantly stronger information model than the first-order query model in previous chapters, which corresponds to being able to evaluate the score just at time 0 ($\nabla \log \pi_0 = \nabla \log \pi$), but we will also allow the evaluations to be inexact.

Despite remaining agnostic to the origin of the score functions, in order to keep our presentation self-contained and to better connect with practice, we discuss the main way in which score functions are obtained in Section 12.2—namely, they are learned from data (samples from π).

- 3 *Discretization.* Finally, our main goal, tackled in Section 12.3, is to control the error incurred by discretization. The key takeaway is that under the stronger information model of accessing the score functions along a diffusion, polynomial sampling complexity bounds can be established with minimal assumptions on π . In particular, we do *not* require π to be log-concave, or to satisfy a functional inequality such as a log-Sobolev inequality.

We now specialize (12.1.1) and (12.1.2) to the case in which the drift $f_t(\cdot)$ is a linear function:

$$dX_t^\rightarrow = -a_t X_t^\rightarrow + g_t dB_t, \quad dX_t^\leftarrow = \{a_{T-t} X_t^\leftarrow + g_{T-t}^2 \nabla \log \pi_{T-t}(X_t^\leftarrow)\} dt + g_{T-t} dB_t. \quad (12.1.3)$$

This choice makes both SDEs easier to simulate. In fact, a simple calculation shows that

$$X_t^\rightarrow | X_0^\rightarrow \sim \text{normal}(A_t X_0^\rightarrow, \sigma_t^2 I_d), \quad (12.1.4)$$

where

$$A_t := \exp\left(-\int_0^t a_s ds\right) \quad \text{and} \quad \sigma_t^2 := \int_0^t g_s^2 \exp\left(-2 \int_s^t a_r dr\right) ds. \quad (12.1.5)$$

12.2 Score Matching and Variants

In this section, we assume that we have access to i.i.d. samples $x^1, \dots, x^n \sim \pi$. The key idea developed here, going back to Hyvärinen (2005); Vincent (2011); Song and Ermon (2019), is that learning score functions from data can be reformulated as a regression objective.

Implicit score matching.

There are multiple variants of score matching, and we start with **implicit score matching**. Here, the goal is to learn $\nabla \log \pi$ from the samples $x^1, \dots, x^n$. A standard approach in statistics is to fix a class $\mathcal{S}$ of candidate score functions $s : \mathbb{R}^d \rightarrow \mathbb{R}^d$ and set up a population-level risk functional $\mathcal{R}$ which is minimized at the true score: $s^\star := \nabla \log \pi = \arg \min_{s \in \mathcal{S}} \mathcal{R}(\pi, s)$. Then, we define our estimator $\hat{s}$ to be the *empirical risk minimizer*, in which we substitute the empirical measure $n^{-1} \sum_{i=1}^n \delta_{x^i}$ for the unknown distribution π : $\hat{s} \in \arg \min_{s \in \mathcal{S}} \mathcal{R}(n^{-1} \sum_{i=1}^n \delta_{x^i}, s)$. In classical theory, the candidate class $\mathcal{S}$ should balance the competing demands of expressivity (necessitating a larger class) and generalization (necessitating a smaller class). In modern practice, we take $\mathcal{S}$ to be a highly expressive class of neural networks, with the hope that generalization follows from other effects such as inductive biases and implicit regularization from the training algorithm.

To see how this works, consider taking the risk functional to be the squared $L^2(\pi)$ norm: $\mathcal{R}(\pi, s) := \mathbb{E}_\pi[\|s - \nabla \log \pi\|^2]$. Although natural, we immediately run into an issue when applying the empirical plug-in approach: the score function of the empirical measure $n^{-1} \sum_{i=1}^n \delta_{x^i}$ is not

well-defined. The idea behind implicit score matching is that we can use the special form of the score function, together with integration by parts, to reformulate the risk $\mathcal{R}$ before applying plug-in. Indeed,

$$\begin{aligned}\mathcal{R}(\pi, \mathbf{s}) &= \int \|\mathbf{s} - \nabla \log \pi\|^2 d\pi = \int \{\|\mathbf{s}\|^2 - 2\langle \mathbf{s}, \nabla \log \pi \rangle\} d\pi + C_\pi \\ &= \int \|\mathbf{s}\|^2 d\pi - 2 \int \langle \mathbf{s}, \nabla \pi \rangle + C_\pi = \underbrace{\int \{\|\mathbf{s}\|^2 + 2 \operatorname{div} \mathbf{s}\} d\pi}_{=:\tilde{\mathcal{R}}(\pi, \mathbf{s})} + C_\pi,\end{aligned}$$

where $C_\pi := \int \|\nabla \log \pi\|^2 d\pi$ is a constant that does not depend on the optimization variable $\mathbf{s}$; in particular, $\arg \min \mathcal{R}(\pi, \cdot) = \arg \min \tilde{\mathcal{R}}(\pi, \cdot)$. However, the latter objective is amenable to plug-in, and yields the implicit score matching estimator $\hat{\mathbf{s}} \in \arg \min_{\mathbf{s} \in \mathcal{S}} n^{-1} \sum_{i=1}^n \{\|\mathbf{s}(x^i)\|^2 + 2 \operatorname{div} \mathbf{s}(x^i)\}$.

In the context of diffusion generative models, our goal is not to learn $\nabla \log \pi$, but rather $\nabla \log \pi_t$ for all $0 < t \leq T$. To apply implicit score matching, observe that given samples from π , we can easily generate samples from π_t via (12.1.4) and (12.1.5).

Implicit score matching applies quite generally, but has the drawback that the training objective requires evaluating the divergence of the candidate score function. Next, we consider a version of score matching which avoids this issue, by explicitly taking advantage of the noising structure.

Denoising score matching.

Here, the setup is as follows: assume that we have samples $x^1, \dots, x^n \sim \pi^X$, but we actually want to learn the score of π^Y , where Y is obtained from X via a known noising process. We now use the following alternative reformulation of the squared error, where C denotes a constant that can change from line from line, and which does not depend on $\mathbf{s}$:

$$\begin{aligned}\int \|\mathbf{s} - \nabla \log \pi^Y\|^2 d\pi^Y &= \int \|\mathbf{s}\|^2 d\pi^Y - 2 \int \langle \mathbf{s}, \nabla \pi^Y \rangle + C \\ &= \int \|\mathbf{s}\|^2 d\pi^Y - 2 \int \left\langle \mathbf{s}(y), \nabla_y \int \pi^{Y|X}(y | x) \pi^X(dx) \right\rangle dy + C \\ &= \int \|\mathbf{s}\|^2 d\pi^Y - 2 \iint \langle \mathbf{s}(y), \nabla_y \pi^{Y|X}(y | x) \rangle \pi^X(dx) dy + C \\ &= \int \|\mathbf{s}\|^2 d\pi^Y - 2 \iint \langle \mathbf{s}(y), \nabla_y \log \pi^{Y|X}(y | x) \rangle \pi^{Y|X}(dy | x) \pi^X(dx) + C \\ &= \int \|\mathbf{s}(y) - \nabla_y \log \pi^{Y|X}(y | x)\|^2 \pi^{Y|X}(dy | x) \pi^X(dx) + C.\end{aligned}$$

Since the noising process $X \rightarrow Y$ is known to us, we can now convert this to an empirical objective. For each sample x^i , generate $y^i \sim \pi^{Y|X}(\cdot | x^i)$. Then, the denoising score estimator is $\hat{\mathbf{s}} \in \arg \min_{\mathbf{s} \in \mathcal{S}} n^{-1} \sum_{i=1}^n \|\mathbf{s}(y^i) - \nabla_y \log \pi^{Y|X}(y^i | x^i)\|^2$. Instead of matching the unconditional score $\nabla \log \pi^Y$, we match the conditional score $\nabla \log \pi^{Y|X}$, which is more tractable.

We can apply this to the diffusion model as follows. Generate $z^1, \dots, z^n \sim \text{normal}(0, I_d)$ independently and independent of $x^1, \dots, x^n$. By (12.1.4) and (12.1.5), the denoising score estimator of $\nabla \log \pi_t$ is given by $\hat{\mathbf{s}}_t \in \arg \min_{\mathbf{s} \in \mathcal{S}} n^{-1} \sum_{i=1}^n \|\mathbf{s}(x_t^i) - (A_t x^i - x_t^i) / \sigma_t^2\|^2$, with $x_t^i := A_t x^i + \sigma_t z^i$. In practice, the scores across time $t \in [0, T]$ are learned jointly through an integrated training objective, but we omit further details here.

What is the goal of generative modelling?

If we step back for a moment, the premise could be confusing: we assume that we have samples $x^1, \dots, x^n \sim \pi$, we use the samples to learn score functions, and then we run the reverse SDE in order to sample from π . Thus, it appears that the goal is to “use samples in order to generate new samples”. But as stated, the premise is absurd: the new samples are not genuinely new, as they are dependent on the previous ones. In fact, in an information-theoretic sense, it is impossible to generate even one fully new (i.e., independent) sample from an existing batch of samples. Otherwise, this would let us amplify our data set for free, contradicting a core tenet of statistics—there is a limited amount of information we can learn from a fixed number of samples.

Note that the tasks “given samples, learn the score functions” and “given access to score function evaluations, generate samples” are each valid in the sense that we can meaningfully search for the best performing methods under the given information structures. It is only the composition of the two that seems, at first, ill-posed.

Although we do not intend to begin a philosophical discourse here, this issue seems significant enough to merit a brief discussion here, with the caveat that the perspective presented here is not the end of the story. We merely claim that it is one logically consistent interpretation.

The perspective we put forward is that the goal of generative modelling is simply distribution learning: given samples from an unknown distribution π , output an approximately correct representation of π . It may help to first consider the case of density estimation, in which the chosen representation is an approximation of the density function p of π (w.r.t. some reference measure μ). In this context, an estimator is a measurable function $\hat{p}$ of the data $x^1, \dots, x^n$, and to explicitly denote this dependence, let us write $\hat{p}_{x^1, \dots, x^n}$. Then, suitable notions of risk could include the integrated mean squared error $\mathbb{E} \int |\hat{p}_{x^1, \dots, x^n} - p|^2 d\mu$ or the total variation error $\frac{1}{2} \mathbb{E} \int |\hat{p}_{x^1, \dots, x^n} - p| d\mu$.

The case of generative modelling is similar, but the chosen representation of π is the *law of the output of a sampler*. A sampler is a randomized algorithm $\mathcal{A}$, which can be thought of as a function of a random seed $U \sim \text{uniform}([0, 1])$. Here, the sampler is learned from data, so we can denote this dependence as $\hat{\mathcal{A}}_{x^1, \dots, x^n}$. If we measure the error in, say, total variation distance, the risk becomes $\mathbb{E} \|\text{law}(\hat{\mathcal{A}}_{x^1, \dots, x^n}(U)) - \pi\|_{\text{TV}}$, or equivalently $\mathbb{E} \|(\hat{\mathcal{A}}_{x^1, \dots, x^n})_{\#} \text{uniform}([0, 1]) - \pi\|_{\text{TV}}$. For a diffusion model, learning a sampler simply amounts to learning the score functions, since the learned score functions (together with a chosen discretization scheme) specify the sampler. We remark that this notion of distribution learning is classical, see [Kearns et al. \(1994\)](#).

Variants.

We pause to describe a few variants of the basic diffusion models we have introduced thus far. Of course, this treatment is far from exhaustive, but serves to convey some of the flexibility in the core design components.

First, we observe that for the purpose of generative modelling, we can freely modify the reverse process, provided that it correctly tracks the marginal laws $(\pi_t)_{t \in [0, T]}$. To see how this is done, write down the Fokker–Planck equation for (12.1.1):

$$\partial_t \pi_t = -\text{div}(\pi_t f_t) + \frac{1}{2} g_t^2 \Delta \pi_t.$$

For the reversal $\pi_t^\leftarrow := \pi_{T-t}$, for any non-negative process $(\tilde{g}_t)_{t \in [0, T]}$,

$$\begin{aligned} \partial_t \pi_t^\leftarrow &= -\partial_t \pi_{T-t} = \operatorname{div}(\pi_{T-t} f_{T-t}) - \frac{1}{2} g_{T-t}^2 \Delta \pi_{T-t} \\ &= -\operatorname{div}\left(\pi_t^\leftarrow \left(-f_{T-t} + \frac{1}{2} (g_{T-t}^2 + \tilde{g}_{T-t}^2) \nabla \log \pi_t^\leftarrow\right)\right) + \frac{1}{2} \tilde{g}_{T-t}^2 \Delta \pi_t^\leftarrow \end{aligned}$$

which corresponds to the SDE

$$dX_t^\leftarrow = \left\{ -f_{T-t}(X_t^\leftarrow) + \frac{1}{2} (g_{T-t}^2 + \tilde{g}_{T-t}^2) \nabla \log \pi_t^\leftarrow(X_t^\leftarrow) \right\} dt + \tilde{g}_{T-t} dB_t.$$

The SDE (12.1.2) corresponds to the choice $\tilde{g}_t = g_t$, but we see that we can make the diffusion coefficient arbitrary as long as we adjust the drift accordingly. In fact, we can even set $\tilde{g}_t = 0$ and we obtain an ODE:

$$\dot{X}_t^\leftarrow = -f_{T-t}(X_t^\leftarrow) + \frac{1}{2} g_{T-t}^2 \nabla \log \pi_t^\leftarrow(X_t^\leftarrow).$$

This is known as the *probability flow ODE* (Song et al., 2021b), and it is equivalent to *denoising diffusion implicit models (DDIMs)* (Song et al., 2021a). For many applications, ODE-based formulations are favoured for the possibility of speeding up generation.

Then came a number of concurrent works essentially sharing the same core observation: rather than restricting ourselves to the time reversal of a forward SDE, we can instead directly define an interpolation path between a reference γ and the target π , and then learn vector fields which transport along the specified path via regression objectives. Frameworks which apply this idea include *flow matching* (Lipman et al., 2023), *rectified flow* (Liu et al., 2023), and *stochastic interpolants* (Albergo and Vanden-Eijnden, 2023). Also, notably, the Schrödinger bridge from Section 3.5 can be interpreted as an entropic variant (De Bortoli et al., 2021; Vargas et al., 2021; Chen et al., 2022a).

Here, we choose to present the framework of stochastic interpolants. Let (x_0, x_T) be drawn from any coupling $\rho \in \mathcal{C}(\gamma, \pi)$. Historically, ρ was taken to be the independent coupling $\gamma \otimes \pi$, but it was soon realized that more general couplings, such as the optimal transport coupling, are also permitted (Pooladian et al., 2023; Albergo et al., 2024; Tong et al., 2024). We define an interpolant $(x_t)_{t \in [0, T]}$, that is, a path that agrees with x_0 and x_T at the endpoints, by specifying the following ingredients for each $t \in [0, T]$. First, an interpolant function $\mathcal{I}_t : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}^d$ such that $\mathcal{I}_0(x_0, x_T) = x_0$ and $\mathcal{I}_T(x_0, x_T) = x_T$; second, a noise variance σ_t^2 , where $\sigma_0^2 = \sigma_T^2 = 0$. Then, a suitable interpolant is given by $x_t := \mathcal{I}_t(x_0, x_T) + \sigma_t z$, where $z \sim \text{normal}(0, I_d)$ is independent of (x_0, x_T) . A simple example to consider is the linear interpolation, $\mathcal{I}_t(x_0, x_T) := \frac{T-t}{T} x_0 + \frac{t}{T} x_T$.

Let $\mu_t := \text{law}(x_t)$. We compute the marginal evolution, as usual, by considering a test function $\varphi : \mathbb{R}^d \rightarrow \mathbb{R}$. Under suitable regularity conditions,

$$\begin{aligned} \partial_t \mathbb{E}_{\mu_t} \varphi &= \partial_t \mathbb{E} \varphi(x_t) = \mathbb{E} \langle \nabla \varphi(x_t), \dot{x}_t \rangle = \mathbb{E} \langle \nabla \varphi(x_t), \dot{\mathcal{I}}_t(x_0, x_T) + \dot{\sigma}_t z \rangle \\ &= \mathbb{E} \langle \nabla \varphi(x_t), \underbrace{\mathbb{E}[\dot{\mathcal{I}}_t(x_0, x_T) + \dot{\sigma}_t z \mid x_t]}_{=: v_t^*(x_t)} \rangle. \end{aligned}$$

It follows that $(\mu_t, v_t^*)_{t \in [0, T]}$ solves the continuity equation (1.3.18). Hence, for the ODE

$$\dot{X}_t = v_t^*(X_t), \quad X_0 \sim \gamma,$$

we have $X_T \sim \pi$. This yields a generative model for sampling from π , provided that we can learn the vector fields $(v_t^*)_{t \in [0, T]}$.

For this, we use the fact that conditional expectations are solutions to regression problems—this is nothing more than their characterization as L^2 projections. Explicitly, for a candidate vector field v_t , we can reformulate the mean squared error as follows:

$$\begin{aligned}\mathbb{E}_{\mu_t}[\|v_t - v_t^*\|^2] &= \mathbb{E}[\|v_t(x_t)\|^2 - 2\langle v_t(x_t), v_t^*(x_t) \rangle] + C \\ &= \mathbb{E}[\|v_t(x_t)\|^2 - 2\langle v_t(x_t), \mathbb{E}[\dot{I}_t(x_0, x_T) + \dot{\sigma}_t z \mid x_t] \rangle] + C \\ &= \mathbb{E}[\|v_t(x_t)\|^2 - 2\langle v_t(x_t), \dot{I}_t(x_0, x_T) + \dot{\sigma}_t z \rangle] + C \\ &= \mathbb{E}[\|v_t(x_t) - \dot{I}_t(x_0, x_T) - \dot{\sigma}_t z\|^2] + C.\end{aligned}$$

This leads to the following training procedure, starting with samples $(x_0^i, x_T^i) \sim \rho$, $i \in [n]$. Generate independent noise variables $z^i \sim \text{normal}(0, I_d)$ and form $x_t^i := \mathcal{I}_t(x_0^i, x_T^i) + \sigma_t z^i$. Then, the estimator for v_t^* is given by $\hat{v}_t \in \arg \min_{v_t \in \mathcal{V}} n^{-1} \sum_{i=1}^n \|v_t(x_t^i) - \dot{I}_t(x_0^i, x_T^i) - \dot{\sigma}_t z^i\|^2$.

In summary, this framework affords flexibility via the choice of coupling ρ , interpolant $\mathcal{I}$, and noise schedule σ . It can also be extended to allow for stochastic generation, although we skip this for brevity. Incidentally, the above reasoning shows that any conditional expectation can be learned via regression, which generalizes the previous derivation of denoising score matching via **Tweedie's identity** (Miyasawa, 1961):

$$\nabla \log \pi^Y(y) = \mathbb{E}[\nabla \log \pi^{Y|X}(Y \mid X) \mid Y = y], \quad X \sim \pi^X, \quad Y \mid X \sim \pi^{Y|X}. \quad (12.2.1)$$

In the subsequent section, we return to the basic diffusion model (12.1.3); in fact, we will further simplify it by fixing the coefficients. This is purely motivated by mathematical simplicity, and does not imply that the original diffusion model formulation is the most important or useful one in practice.

12.3 Discretization Analysis

We now turn toward discretization analysis, and for maximum simplicity, we take $a_t = 0$ and $g_t = 1$ in (12.1.3). Thus, we end up with the processes

$$dX_t^\rightarrow = dB_t, \quad dX_t^\leftarrow = \nabla \log \pi_{T-t}(X_t^\leftarrow) dt + dB_t, \quad \pi_t = \pi * \text{normal}(0, tI_d).$$

This is sometimes known as the “variance-exploding” (VE) parametrization.

In the following analysis, a special role will be played by the quantities

$$D_t^*(x_t) := \mathbb{E}[X_0^\rightarrow \mid X_t^\rightarrow = x_t], \quad \Sigma_t(x_t) := \text{cov}(X_0^\rightarrow \mid X_t^\rightarrow = x_t).$$

Here, we use the symbol D_t^* because it is the optimal denoiser. By Tweedie's identity (12.2.1), we can rewrite the reverse process as

$$dX_t^\leftarrow = \frac{D_{T-t}^*(X_t^\leftarrow) - X_t^\leftarrow}{T-t} dt + dB_t, \quad X_0^\leftarrow \sim \pi_T. \quad (12.3.1)$$

Given an approximate score function $\hat{\mathbf{s}}_t$, let $\hat{D}_t$ denote the corresponding denoiser:

$$\hat{D}_t(x_t) := x_t + t \hat{\mathbf{s}}_t(x_t) \iff \hat{\mathbf{s}}_t(x_t) = \frac{\hat{D}_t(x_t) - x_t}{t}.$$

For the best results, we will need a *time-varying* sequence of step sizes $(h_n)_{n=0,1,\dots,N-1}$. Let $\tau_0 := 0$ and $\tau_n := \sum_{k=0}^{n-1} h_k$, and for $\tau_n \leq t < \tau_{n+1}$ we write $t_- := \tau_n$. We consider the discretization which

freezes only the term $\mathbf{D}_{T-t}^*(X_t^\leftarrow)$ in (12.3.1). Also, for large T , we expect that $\pi_T \approx \text{normal}(0, TI_d)$, which will serve as our initialization. Hence, our overall algorithm is:

$$d\hat{X}_t^\leftarrow = \frac{\widehat{\mathbf{D}}_{T-t}(\hat{X}_t^\leftarrow) - \hat{X}_t^\leftarrow}{T-t} dt + dB_t, \quad \hat{X}_0^\leftarrow \sim \gamma := \text{normal}(0, TI_d). \quad (\text{DDPM})$$

We refer to this as the *denoising diffusion probabilistic model* (DDPM), as is common in the literature. Using a bit of stochastic calculus, this leads to the more explicit update

$$\hat{X}_{\tau_{n+1}}^\leftarrow = \frac{T - \tau_{n+1}}{T - \tau_n} \hat{X}_{\tau_n}^\leftarrow + \frac{h_n}{T - \tau_n} \widehat{\mathbf{D}}_{T-\tau_n}(\hat{X}_{\tau_n}^\leftarrow) + \text{normal}\left(0, \frac{T - \tau_{n+1}}{T - \tau_n} h_n I_d\right).$$

At first glance, this may seem like a rather complicated choice of discretization, but as we shall see, it engenders a particularly elegant and refined analysis.

A primary consideration here is to perform the analysis under minimal assumptions on π . Indeed, in practical applications of generative modelling, it is unreasonable to assume that the target distribution π satisfies conditions such as smoothness or log-concavity. In fact, we should not even assume that π admits a Lebesgue density: if π is a distribution over “natural” images, then it is often suggested that π could be supported on a low-dimensional manifold (the “manifold hypothesis”). This may seem worrying, but it will not pose an issue, because DDPM does not aim to sample directly from π . At the terminal time, if all goes well, at best we can expect $\text{law}(\hat{X}_{\tau_N}^\leftarrow) \approx \pi_{T-\tau_N} = \pi_\delta$, where $\delta := T - \tau_N$ is called the *early stopping* parameter. Thus, our goal will be to show that we can sample from π_δ with an iteration complexity that scales mildly with $1/\delta$ —in fact, we will obtain logarithmic dependence—with the idea that if we set δ to be extremely small, then π and π_δ are practically indistinguishable. For example, it is easy to see that $W_2(\pi, \pi_\delta) \leq \sqrt{\delta d}$.

With the above discussion in mind, the *only* assumption that we will impose at first is that π merely has a finite second moment. Then, we will show that the bounds can be refined if π really does have low-dimensional structure, consistent with the manifold hypothesis.

The other important issue is the assumption quantifying the agreement between the approximate score $\widehat{\mathbf{s}}_t$ and the true score $\mathbf{s}_t$. Our analysis assumes a bound on the L^2 norm $\|\widehat{\mathbf{s}}_t - \mathbf{s}_t\|_{L^2(\pi_t)}$. This is a desirable feature of the analysis, since it exactly corresponds to the excess population-level risk in the earlier derivations of the score matching objectives. Thus, if the score functions are learned from data via those training objectives, we can hope (optimistically) that this assumption on the approximate score functions is indeed satisfied.

Without further ado, we present the main theorem.

Theorem 12.3.2 (Analysis of DDPM). *Let the step sizes be chosen so that $h_n \leq \bar{h} (T - \tau_n)$ for some $\bar{h} \leq 1/2$. Then,*

$$\text{KL}(\pi_\delta \parallel \hat{\pi}_{\tau_N}) \leq \frac{m_2^2}{2T} + 4\varepsilon_{\text{sc}}^2 + 2\bar{h} \mathfrak{D}_{\delta,T}(\pi),$$

where the three terms correspond to initialization error, score error, and discretization error respectively. Here,

$$m_2^2 := \mathbb{E}_\pi[\|\cdot\|^2], \quad \varepsilon_{\text{sc}}^2 := \sum_{n=0}^{N-1} h_n \mathbb{E}_{\pi_{T-\tau_n}}[\|\hat{\mathbf{S}}_{T-\tau_n} - \mathbf{s}_{T-\tau_n}^*\|^2],$$

and

$$\mathfrak{D}_{\delta,T}(\pi) := \frac{\mathbb{E}_{\pi_T} \text{tr} \Sigma_T}{T} + \int_\delta^T \frac{\mathbb{E}_{\pi_t} \text{tr} \Sigma_t}{t^2} dt.$$

In particular, if we take $\tau_n = (1 - (1 - \bar{h})^n)T$, then for $T \asymp m_2^2/\varepsilon^2$ and $\bar{h} \asymp \varepsilon^2/\mathfrak{D}_{\delta,T}(\pi)$, we obtain $\text{KL}(\pi_\delta \parallel \hat{\pi}_{\tau_N}) \lesssim \varepsilon^2 + \varepsilon_{\text{sc}}^2$ with

$$N \asymp \frac{\mathfrak{D}_{\delta,T}(\pi)}{\varepsilon^2} \log \frac{m_2^2}{\delta \varepsilon^2} \quad \text{score evaluations}.$$

Note that we can only sample up to error $\varepsilon_{\text{sc}}^2$ at best, so for the success of DDPM, it is crucial that we have accurate score functions.

Since the result depends on π via $\mathfrak{D}_{\delta,T}(\pi)$, our next task is to interpret this quantity. We provide a first bound that applies to any distribution π with finite second moment.

Lemma 12.3.3. *For any probability measure π with finite second moment,*

$$\mathfrak{D}_{\delta,T}(\pi) \leq d (1 + \log(T/\delta)).$$

Proof For any $t > 0$,

$$\mathbb{E}_{\pi_t} \text{tr} \Sigma_t = \mathbb{E}[\|X_0^\rightarrow - \mathbb{E}[X_0^\rightarrow \mid X_t^\rightarrow]\|^2] \leq \mathbb{E}[\|X_0^\rightarrow - X_t^\rightarrow\|^2] = dt. \quad (12.3.4)$$

Substitute this into the definition of $\mathfrak{D}_{\delta,T}(\pi)$. $\square$

This readily yields that, provided that we are not bottlenecked by the score error, the iteration complexity for DDPM to sample from π_δ is $O((d/\varepsilon^2) \log^2(m_2^2/\delta \varepsilon^2))$.

However, we can say more.

Lemma 12.3.5. *For any probability measure π with finite second moment and any $T \geq \delta$,*

$$\mathfrak{D}_{\delta,T}(\pi) \leq 4 \{ \mathcal{H}(\text{normal}(0, \delta I_d)) - \mathcal{H}(\pi_\delta) \},$$

where $\mathcal{H}(\cdot)$ is the entropy functional.

Proof First, since $t \mapsto \mathbb{E}_{\pi_t} \text{tr} \Sigma_t$ is increasing (in fact, we shall obtain a formula for its derivative

in (12.3.10) below),

$$\frac{\mathbb{E}_{\pi_T} \operatorname{tr} \Sigma_T}{T} = \mathbb{E}_{\pi_T} \operatorname{tr} \Sigma_T \int_T^\infty \frac{1}{t^2} dt \leq \int_T^\infty \frac{\mathbb{E}_{\pi_t} \operatorname{tr} \Sigma_t}{t^2} dt.$$

Therefore,

$$\mathfrak{D}_{\delta,T}(\pi) \leq 2 \int_\delta^\infty \frac{\mathbb{E}_{\pi_t} \operatorname{tr} \Sigma_t}{t^2} dt.$$

Next, by orthogonality,

$$\mathbb{E}_{\pi_t} \operatorname{tr} \Sigma_t = \mathbb{E}[\|X_0^\rightarrow - X_t^\rightarrow\|^2] - \mathbb{E}[\|D_t^\star(X_t^\rightarrow) - X_t^\rightarrow\|^2] = dt - t^2 \mathbb{E}_{\pi_t}[\|\mathbf{s}_t^\star\|^2].$$

Note that this refines (12.3.4). This yields

$$\mathfrak{D}_{\delta,T}(\pi) \leq 2 \int_\delta^\infty \left(\frac{d}{t} - \mathbb{E}_{\pi_t}[\|\mathbf{s}_t^\star\|^2] \right) dt = 2 \lim_{T \rightarrow \infty} \int_\delta^T \left(\frac{d}{t} - \mathbb{E}_{\pi_t}[\|\mathbf{s}_t^\star\|^2] \right) dt.$$

The idea now is to use the fact that $\partial_t \mathcal{H}(\pi_t) = -\frac{1}{2} \mathbb{E}_{\pi_t}[\|\mathbf{s}_t^\star\|^2]$, where $\mathcal{H}(\cdot)$ is the entropy functional. Moreover, if $\gamma_t := \operatorname{normal}(0, tI_d)$, then we can recognize $d/t = \mathbb{E}_{\gamma_t}[\|\nabla \log \gamma_t\|^2]$. Hence,

$$\mathfrak{D}_{\delta,T}(\pi) \leq 4 \lim_{T \rightarrow \infty} \{ -\mathcal{H}(\gamma_T) + \mathcal{H}(\gamma_\delta) + \mathcal{H}(\pi_T) - \mathcal{H}(\pi_\delta) \}.$$

We claim that $\lim_{T \rightarrow \infty} \{ \mathcal{H}(\pi_T) - \mathcal{H}(\gamma_T) \} = 0$; this can be seen, e.g., via

$$|\mathcal{H}(\pi_T) - \mathcal{H}(\gamma_T)| = \left| \operatorname{KL}(\pi_T \parallel \gamma_T) + \frac{d}{2} - \frac{1}{2T} \int \|\cdot\|^2 d\pi_T \right| = \left| \operatorname{KL}(\pi_T \parallel \gamma_T) - \frac{m_2^2}{2T} \right| \leq \frac{m_2^2}{2T},$$

by (12.3.8) below. This completes the proof. $\square$

One way to bound the entropy is via a covering number condition.

Lemma 12.3.6. *Suppose that π is compactly supported, and let $\mathcal{N}(\operatorname{supp} \pi, r)$ denote the covering number of the support of π at scale r . Then,*

$$\mathcal{H}(\operatorname{normal}(0, \delta I_d)) - \mathcal{H}(\pi_\delta) \leq \log \mathcal{N}(\operatorname{supp} \pi, \sqrt{\delta}) + \frac{1}{2}.$$

Proof In our convention, the entropy functional $\mathcal{H}$ is the *negative* of the usual (Shannon) differential entropy. Since the following proof is best phrased in the language of information theory, let us temporarily write $\mathcal{H} := -\mathcal{H}$.

Let $B_1, \dots, B_N$ be a covering of $\operatorname{supp} \pi$ by balls of radius r , where $N := \mathcal{N}(\operatorname{supp} \pi, r)$. Let J be a random variable denoting the index of a ball in which $X_0^\rightarrow$ lies; thus, $X_0^\rightarrow \in B_J$. Then,

$$\mathcal{H}(X_\delta^\rightarrow) - \mathcal{H}(X_\delta^\rightarrow \mid J) = \mathcal{I}(J; X_\delta^\rightarrow) = \mathcal{H}(J) - \mathcal{H}(J \mid X_\delta^\rightarrow),$$

where $\mathcal{H}$ and $\mathcal{I}$ are the discrete Shannon entropy and mutual information respectively; see Exercise 9.1. This implies that

$$\mathcal{H}(X_\delta^\rightarrow) = \mathcal{H}(J) - \mathcal{H}(J \mid X_\delta^\rightarrow) + \mathcal{H}(X_\delta^\rightarrow \mid J) \leq \log N + \max_{j \in [N]} \mathcal{H}(X_\delta^\rightarrow \mid J = j).$$

Next, recall that Gaussians maximize the entropy under a fixed covariance constraint, and that the entropy of a Gaussian with covariance Σ is $\frac{1}{2} \log((2\pi e)^d \det \Sigma)$. Hence, by the AM–GM inequality,

$$\begin{aligned} \mathcal{H}(X_\delta^\rightarrow \mid J = j) &\leq \frac{1}{2} \log((2\pi e)^d \det \text{cov}(X_\delta^\rightarrow \mid J = j)) \leq \frac{d}{2} \log \frac{2\pi e \text{tr} \text{cov}(X_\delta^\rightarrow \mid J = j)}{d} \\ &\leq \frac{d}{2} \log \frac{2\pi e (r^2 + \delta d)}{d} = \mathcal{H}(\text{normal}(0, \delta I_d)) + \frac{d}{2} \log\left(1 + \frac{r^2}{\delta d}\right). \end{aligned}$$

Choose $r = \sqrt{\delta}$ and note that $\frac{d}{2} \log(1 + \frac{1}{d}) \leq \frac{1}{2}$. $\square$

Combining the lemmas shows that the iteration complexity of DDPM is $O((k_\delta(\pi)/\varepsilon^2) \log(m_2^2/\delta\varepsilon^2))$, where $k_\delta(\pi) := \log \mathcal{N}(\text{supp } \pi, \sqrt{\delta})$. The quantity $k_\delta(\pi)$ is a measure of the *intrinsic dimensionality* of π —for example, if π is supported on a ball of radius R in a k -dimensional subspace, then $k_\delta(\pi) \asymp k \log(R/\sqrt{\delta})$. However, this notion is rather flexible and allows for π to be supported close (up to scale $\sqrt{\delta}$) to a k -dimensional manifold, or a finite union thereof, etc.

Remark 12.3.7. How do we interpret closeness to the early stopped measure? One answer is that by using $W_2(\pi, \pi_\delta) \leq \sqrt{\delta d}$, we can state a guarantee in a common metric that is weaker than both W_2 and $\sqrt{\text{KL}}$. As an example, consider the bounded Lipschitz metric

$$\|\mu - \nu\|_{\text{BL}} = \sup \left\{ \int f d(\mu - \nu) \mid |f| \leq 1, f \text{ is 1-Lipschitz} \right\}.$$

Then, $\|\mu - \nu\|_{\text{BL}} \leq \|\mu - \nu\|_{\text{TV}} \wedge W_1(\mu, \nu)$. This metric is known to metrize weak convergence of probability measures. The guarantee of Theorem 12.3.2 then implies $\|\hat{\pi}_{\tau_N} - \pi\|_{\text{BL}} \lesssim \varepsilon_{\text{sc}} + \varepsilon + \sqrt{\delta d}$.

If we further know that π is β -log-smooth, then

$$\begin{aligned} \partial_t \text{KL}(\pi_t \parallel \pi) &= -\frac{1}{2} \mathbb{E}_{\pi_t} \left\langle \nabla \log \frac{\pi_t}{\pi}, \nabla \log \pi_t \right\rangle \leq \frac{1}{4} \mathbb{E}_{\pi_t} [\|\nabla \log \pi\|^2 + \|\nabla \log \pi_t\|^2] \\ &\leq \frac{1}{4} (2 \mathbb{E}_{\pi} [\|\nabla \log \pi\|^2] + 2\beta^2 W_2^2(\pi_t, \pi) + \mathbb{E}_{\pi_t} [\|\nabla \log \pi_t\|^2]) \\ &\leq \frac{3}{4} \mathbb{E}_{\pi} [\|\nabla \log \pi\|^2] + \frac{\beta^2 dt}{2}, \end{aligned}$$

where we used the fact that the Fisher information is decreasing along the heat flow; see (2.2.16). If we further use $\mathbb{E}_{\pi} [\|\nabla \log \pi\|^2] \leq \beta d$ (Lemma 4.0.1) and integrate,

$$\text{KL}(\pi_\delta \parallel \pi) \leq \frac{3\beta\delta d + \beta^2\delta^2 d}{4}.$$

This could be used to obtain bounds on $\|\hat{\pi}_{\tau_N} - \pi\|_{\text{TV}}$.

We now present the proof of the main theorem.

Proof of Theorem 12.3.2 We apply the Girsanov method from Section 4.4. Let $\hat{\mathbf{P}}, \mathbf{P}$ denote the path measures corresponding to DDPM and (12.3.1) respectively. Also, let $\hat{\mathbf{P}}_x, \mathbf{P}_x$ denote the corresponding measures conditioned on starting at x . By the chain rule for the KL divergence (Lemma 1.5.8) and Girsanov's theorem (Theorem 3.2.8),

$$\begin{aligned} \text{KL}(\mathbf{P} \parallel \hat{\mathbf{P}}) &= \text{KL}(\pi_T \parallel \gamma) + \int \text{KL}(\mathbf{P}_x \parallel \hat{\mathbf{P}}_x) \pi_T(dx) \\ &= \text{KL}(\pi_T \parallel \gamma) + \frac{1}{2} \mathbb{E} \int_0^{\tau_N} \frac{\|\widehat{\mathbf{D}}_{T-t}(X_t) - \mathbf{D}_{T-t}^*(X_t)\|^2}{(T-t)^2} dt, \end{aligned}$$

where $(X_t)_{t \in [0, T]}$ denotes the ideal reverse process (12.3.1); to simplify the notation, we have dropped the $\leftarrow$ superscript. The first term captures the effect of initialization. By the joint convexity of the KL divergence, it can be bounded as follows:

$$\begin{aligned} \text{KL}(\pi_T \parallel \gamma) &= \text{KL}(\pi * \text{normal}(0, TI_d) \parallel \text{normal}(0, TI_d)) \\ &\leq \int \text{KL}(\text{normal}(x, TI_d) \parallel \text{normal}(0, TI_d)) \pi(dx) = \frac{1}{2T} \int \|x\|^2 \pi(dx) = \frac{m_2^2}{2T}. \end{aligned} \quad (12.3.8)$$

Alternatively, this can be seen as a consequence of the log-Harnack inequality; see [Exercise 2.17](#), [Exercise 3.4](#), and [Exercise 4.7](#).

For the second term, let us separate out the effect of the score error:

$$\begin{aligned} &\frac{1}{2} \mathbb{E} \int_0^{\tau_N} \frac{\|\widehat{D}_{T-t_-}(X_{t_-}) - D_{T-t}^*(X_t)\|^2}{(T-t)^2} dt \\ &\leq \mathbb{E} \int_0^{\tau_N} \frac{\|\widehat{D}_{T-t_-}(X_{t_-}) - D_{T-t_-}^*(X_{t_-})\|^2 + \|D_{T-t}^*(X_t) - D_{T-t_-}^*(X_{t_-})\|^2}{(T-t)^2} dt \\ &= \int_0^{\tau_N} \left(\frac{T-t_-}{T-t} \right)^2 \mathbb{E}[\|\widehat{s}_{T-t_-}(X_{t_-}) - s_{T-t_-}^*(X_{t_-})\|^2] dt + \mathbb{E} \int_0^{\tau_N} \frac{\|D_{T-t}^*(X_t) - D_{T-t_-}^*(X_{t_-})\|^2}{(T-t)^2} dt \\ &\leq 4 \int_0^{\tau_N} \mathbb{E}[\|\widehat{s}_{T-t_-}(X_{t_-}) - s_{T-t_-}^*(X_{t_-})\|^2] dt + \mathbb{E} \int_0^{\tau_N} \frac{\|D_{T-t}^*(X_t) - D_{T-t_-}^*(X_{t_-})\|^2}{(T-t)^2} dt, \end{aligned}$$

where we use the assumption on the step size to conclude that $T - t_- \leq 2(T - t)$. The first term is identified as $4\varepsilon_{\text{sc}}^2$. Observe carefully that by bounding $\text{KL}(\mathbf{P} \parallel \widehat{\mathbf{P}})$ with this order of arguments, the expectations are taken with respect to the ideal process (12.3.1). This is what allows us to assume only that the score errors are bounded in L^2 with respect to the true marginals. This also plays an important role in the analysis of the last term, which captures the error incurred from discretization.

A naïve approach would be to assume (or prove) that $(x, t) \mapsto D_t^*(x)$ is Lipschitz in space and time, but it turns out that we can achieve a significantly more refined bound by applying Itô's formula directly to $D_{T-t}^*(X_t)$. This calculation was essentially carried out in [Exercise 3.5](#), and it implies

$$ds_{T-t}^*(X_t) = \nabla s_{T-t}^*(X_t) dB_t.$$

In terms of D^* , we have

$$dD_{T-t}^*(X_t) = \nabla D_{T-t}^*(X_t) dB_t. \quad (12.3.9)$$

We invite the interested reader to verify these identities as an exercise.

Hence, by the Itô isometry (1.1.9),

$$\mathbb{E}[\|D_{T-t}^*(X_t) - D_{T-t_-}^*(X_{t_-})\|^2] = \mathbb{E}\left[\left\|\int_{t_-}^t \nabla D_{T-s}^*(X_s) dB_s\right\|^2\right] = \int_{t_-}^t \mathbb{E}[\|\nabla D_{T-s}^*(X_s)\|_{\text{HS}}^2] ds.$$

Plugging this back in,

$$\begin{aligned} \mathbb{E} \int_0^{\tau_N} \frac{\|\mathbf{D}_{T-t}^*(X_t) - \mathbf{D}_{T-t_-}^*(X_{t_-})\|^2}{(T-t)^2} dt &= \int_0^{\tau_N} \frac{1}{(T-t)^2} \int_{t_-}^t \mathbb{E}[\|\nabla \mathbf{D}_{T-s}^*(X_s)\|_{\text{HS}}^2] ds dt \\ &= \int_0^{\tau_N} \frac{s_+ - s}{(T-s)(T-s_+)} \mathbb{E}[\|\nabla \mathbf{D}_{T-s}^*(X_s)\|_{\text{HS}}^2] ds, \end{aligned}$$

where, for $\tau_n \leq s < \tau_{n+1}$, we write $s_+ := \tau_{n+1}$. For such s , by our assumption on the step sizes, $s_+ - s \leq h_n \leq \bar{h} (T - \tau_n) \leq 2\bar{h} (T - \tau_{n+1}) = 2\bar{h} (T - s_+)$, so

$$\mathbb{E} \int_0^{\tau_N} \frac{\|\mathbf{D}_{T-t}^*(X_t) - \mathbf{D}_{T-t_-}^*(X_{t_-})\|^2}{(T-t)^2} dt \leq 2\bar{h} \int_0^{\tau_N} \frac{\mathbb{E}[\|\nabla \mathbf{D}_{T-s}^*(X_s)\|_{\text{HS}}^2]}{T-s} ds.$$

On the other hand, (12.3.9) formally shows that

$$\mathbf{D}_{T-t}^*(X_t) - X_T = \mathbf{D}_{T-t}^*(X_t) - \mathbf{D}_0^*(X_T) = - \int_t^T \nabla \mathbf{D}_{T-s}^*(X_s) dB_s$$

and hence

$$\mathbb{E}[\|\mathbf{D}_{T-t}^*(X_t) - X_T\|^2] = \int_t^T \mathbb{E}[\|\nabla \mathbf{D}_{T-s}^*(X_s)\|_{\text{HS}}^2] ds.$$

In other words,

$$\begin{aligned} \mathbb{E}[\|\nabla \mathbf{D}_{T-s}^*(X_s)\|_{\text{HS}}^2] &= -\partial_s \mathbb{E}[\|\mathbf{D}_{T-s}^*(X_s) - X_T\|^2] = -\partial_s \mathbb{E}[\|\mathbf{D}_{T-s}^*(X_{T-s}^\rightarrow) - X_0^\rightarrow\|^2] \\ &= -\partial_s \mathbb{E}_{\pi_{T-s}} \text{tr} \Sigma_{T-s} = \partial_t \Big|_{t=T-s} \mathbb{E}_{\pi_t} \text{tr} \Sigma_t. \end{aligned} \quad (12.3.10)$$

Integration by parts now yields

$$\begin{aligned} \int_0^{\tau_N} \frac{\mathbb{E}[\|\nabla \mathbf{D}_{T-s}^*(X_s)\|_{\text{HS}}^2]}{T-s} ds &= \int_0^{\tau_N} \frac{\partial_t \Big|_{t=T-s} \mathbb{E}_{\pi_t} \text{tr} \Sigma_t}{T-s} ds = \int_\delta^T \frac{\partial_t \mathbb{E}_{\pi_t} \text{tr} \Sigma_t}{t} dt \\ &= \frac{\mathbb{E}_{\pi_T} \text{tr} \Sigma_T}{T} - \frac{\mathbb{E}_{\pi_\delta} \text{tr} \Sigma_\delta}{\delta} + \int_\delta^T \frac{\mathbb{E}_{\pi_t} \text{tr} \Sigma_t}{t^2} dt. \end{aligned}$$

Dropping the second term above and gathering together the remaining terms completes the proof. $\square$

Finally, we close this section with some remarks on further improvements. The works of [Li et al. \(2024b\)](#); [Li and Yan \(2025\)](#); [Liang et al. \(2025\)](#) showed improved iteration bounds of $\tilde{O}(d/\varepsilon)$ under minimal assumptions, and $\tilde{O}(k(\pi)/\varepsilon)$ where $k(\pi)$ again denotes the intrinsic dimension of π (at a certain scale). These bounds hold in total variation distance, and the $\tilde{O}(d/\varepsilon)$ rate was later shown to hold in KL divergence in [Jain and Zhang \(2026\)](#).

That there is room for improvement could be predicted by the fact that, by (12.3.9), the drift is a martingale; thus, the weak error, in the sense of [Lemma 5.1.2](#), morally has zero contribution from the discretization error (only a contribution from the score error). However, there is currently no proof of the refined rates along these lines.

It is even possible to obtain high-accuracy guarantees for diffusion models—again, provided that score errors are not the bottleneck. In [Chen et al. \(2026b\)](#), it was shown that an iteration complexity of $O(k(\pi) \log^3(m^2/\delta\varepsilon^2))$ can be achieved via approximate rejection sampling. High-accuracy and dimension-free guarantees were also shown in [Gatmiry et al. \(2026\)](#) for Gaussian mixtures.

Recently, [Xun and Price \(2026\)](#) gave a lower bound in which $\tilde{\Omega}(\sqrt{d})$ queries are needed for diffusion sampling. Here too, fundamental questions regarding the complexity of sampling remain open.

Bibliographical Notes

The literature on generative models is vast and growing, so there will inevitably be important omissions. Some of the main references were already cited in the main text above.

Early works on the discretization analysis of diffusion models include [De Bortoli et al. \(2021\)](#); [Block et al. \(2022\)](#); [De Bortoli \(2022\)](#); [Lee et al. \(2022\)](#); [Liu et al. \(2022\)](#); [Pidstrigach \(2022\)](#). These works either assumed strong assumptions on the data (e.g., a log-Sobolev inequality) or on the score errors (e.g., L^∞ bounds), or did not obtain polynomial iteration complexities. The first works to resolve these issues were [Chen et al. \(2023c\)](#); [Lee et al. \(2023\)](#). The main results for these two works hold under the assumption that the true scores $\{\mathbf{s}_t^*\}_{t \in [0, T]}$ are Lipschitz. When combined with early stopping, this implies polynomial iteration complexity bounds for data distributions which are compactly supported, but the quantitative bounds are quite poor. The bounds were later refined in [Chen et al. \(2023a\)](#), and then in [Benton et al. \(2024\)](#); [Conforti et al. \(2025\)](#) which achieved $\tilde{O}(d/\varepsilon^2)$ complexity bounds under minimal data assumptions.

Then, many further works showed that the iteration complexity can be made to scale with the intrinsic dimensionality: [Li and Yan \(2025\)](#); [Liang et al. \(2025\)](#); [Potapchik et al. \(2025\)](#); [Huang et al. \(2026\)](#). Moreover, when π is a Gaussian mixture, the iteration bounds can be made dimension-free ([Li and Cai, 2025](#); [Chen et al., 2026b](#); [Gatmiry et al., 2026](#)).

Another line of works focuses on improved dependence on d or $1/\varepsilon$ via stronger smoothness assumptions on the score function and/or score errors, possibly by incorporating ODE-based samplers. We remark that the analysis of ODE-based samplers is challenging in its own right, as Girsanov's theorem is no longer available, and the story here is perhaps not yet settled. We do not cover these approaches in this chapter, as we focus on the setting with minimal data assumptions. See, however, [Chen et al. \(2023b\)](#); [Li et al. \(2024a\)](#); [Gupta et al. \(2025\)](#); [Li and Cai \(2025\)](#); [Li and Jiao \(2025\)](#); [Li et al. \(2025b\)](#); [Liu et al. \(2025\)](#); [Zhang et al. \(2025\)](#); [Chen et al. \(2026b\)](#).

Diffusion models can also be sped up via parallelization ([Shih et al., 2023](#); [Kandasamy and Nagaraj, 2024](#); [Gupta et al., 2025](#); [Zhou and Sugiyama, 2025b](#)).

The calculation in [Lemma 12.3.5](#), in the language of information theory, is known as the *I-MMSE identity* ([Guo et al., 2005](#)) as it relates the derivative of the mutual information with the minimum mean squared error of a reconstruction problem. Indeed, $\mathbb{E}_{\pi_t} \text{tr} \Sigma_t = \mathbb{E}[\|X_0^\rightarrow - D_t^*(X_t^\rightarrow)\|^2]$ is the Bayes risk of recovering $X_0^\rightarrow$ given the noisy observation $X_t^\rightarrow$, and since $\mathcal{H}(\text{normal}(0, \delta I_d)) - \mathcal{H}(\pi_\delta) = \mathcal{I}(X_0^\rightarrow; X_\delta^\rightarrow)$, the calculation therein shows that $\mathcal{I}(X_0^\rightarrow; X_\delta^\rightarrow) = \frac{1}{2} \int_\delta^\infty t^{-2} \mathbb{E}_{\pi_t} \text{tr} \Sigma_t dt$, or $\partial_t \mathcal{I}(X_0^\rightarrow; X_t^\rightarrow) = -\frac{1}{2t^2} \mathbb{E}_{\pi_t} \text{tr} \Sigma_t$. Related information-theoretic perspectives include [Kong et al. \(2023\)](#); [Jeon et al. \(2025\)](#); [Reeves and Pfister \(2025\)](#); [Aghapour and Bayraktar \(2026\)](#); [Aghapour et al. \(2026\)](#); [Akhtar et al. \(2026\)](#).

References

- Agarwal, Alekh, Bartlett, Peter L., Ravikumar, Pradeep, and Wainwright, Martin J. 2012. Information-theoretic lower bounds on the oracle complexity of stochastic convex optimization. *IEEE Trans. Inform. Theory*, **58**(5), 3235–3249.
- Aghapour, Ahmad, and Bayraktar, Erhan. 2026. When diffusion model can ignore dimension: an entropy-based theory. *arXiv preprint 2605.07969*.
- Aghapour, Ahmad, Bayraktar, Erhan, and Zhang, Ziqing. 2026. Entropy-based dimension-free convergence and loss-adaptive schedules for diffusion models. *arXiv preprint 2601.21943*.
- Ahn, Kwangjun, and Chewi, Sinho. 2021. Efficient constrained sampling via the mirror-Langevin algorithm. Pages 28405–28418 of: Ranzato, M., Beygelzimer, A., Nguyen, K., Liang, P. S., Vaughan, J. W., and Dauphin, Y. (eds), *Advances in Neural Information Processing Systems*, vol. 34. Curran Associates, Inc.
- Akhtar, Md S., El Gadarri, Aymane, Farias, Vivek F., and Jozefiak, Adam D. 2026. The value of covariance matching in Gaussian DDPMs and the Lanczos sampler. *arXiv preprint 2605.22723*.
- Akyildiz, O. Deniz, and Sabanis, Sotirios. 2024. Nonasymptotic analysis of stochastic gradient Hamiltonian Monte Carlo under local conditions for nonconvex optimization. *Journal of Machine Learning Research*, **25**(113), 1–34.
- Albergo, Michael S., and Vanden-Eijnden, Eric. 2023. Building normalizing flows with stochastic interpolants. In: *The Eleventh International Conference on Learning Representations*.
- Albergo, Michael S., Goldstein, Mark, Boffi, Nicholas M., Ranganath, Rajesh, and Vanden-Eijnden, Eric. 2024. Stochastic interpolants with data-dependent couplings. Pages 921–937 of: Salakhutdinov, Ruslan, Kolter, Zico, Heller, Katherine, Weller, Adrian, Oliver, Nuria, Scarlett, Jonathan, and Berkenkamp, Felix (eds), *Proceedings of the 41st International Conference on Machine Learning*. Proceedings of Machine Learning Research, vol. 235. PMLR.
- Albritton, Dallas, Armstrong, Scott, Mourrat, Jean-Christophe, and Novack, Matthew. 2024. Variational methods for the kinetic Fokker–Planck equation. *Anal. PDE*, **17**(6), 1953–2010.
- Alonso-Gutiérrez, David, and Bastero, Jesús. 2015. *Approaching the Kannan-Lovász-Simonovits and variance conjectures*. Lecture Notes in Mathematics, vol. 2131. Springer, Cham.
- Alquier, Pierre. 2024. User-friendly introduction to PAC–Bayes bounds. *Foundations and Trends® in Machine Learning*, **17**(2), 174–303.
- Altschuler, Jason M., and Chewi, Sinho. 2024a. Faster high-accuracy log-concave sampling via algorithmic warm starts. *J. ACM*, **71**(3), Art. 24, 55.
- Altschuler, Jason M., and Chewi, Sinho. 2024b. Shifted composition I: Harnack and reverse transport inequalities. *IEEE Transactions on Information Theory*, 1–1.
- Altschuler, Jason M., and Chewi, Sinho. 2026a. Shifted composition II: shift Harnack inequalities and curvature upper bounds. *IEEE Transactions on Information Theory*, **72**(1), 42–57.
- Altschuler, Jason M., and Chewi, Sinho. 2026b. Shifted composition III: local error framework for KL divergence. *Found. Comput. Math.*

- Altschuler, Jason M., and Talwar, Kunal. 2023. Resolving the mixing time of the Langevin algorithm to its stationary distribution for log-concave sampling. Pages 2509–2510 of: Neu, Gergely, and Rosasco, Lorenzo (eds), *Proceedings of Thirty Sixth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 195. PMLR.
- Altschuler, Jason M., Chewi, Sinho, and Zhang, Matthew S. 2025. Shifted composition IV: underdamped Langevin and numerical discretizations with partial acceleration. *arXiv preprint 2506.23062*.
- Ambrosio, Luigi, Gigli, Nicola, and Savaré, Giuseppe. 2008. *Gradient flows in metric spaces and in the space of probability measures*. Second edn. Lectures in Mathematics ETH Zürich. Birkhäuser Verlag, Basel.
- Andrieu, Christophe, Durmus, Alain, Nüsken, Nikolas, and Roussel, Julien. 2021a. Hypocoercivity of piecewise deterministic Markov process-Monte Carlo. *Ann. Appl. Probab.*, **31**(5), 2478–2517.
- Andrieu, Christophe, Dobson, Paul, and Wang, Andi Q. 2021b. Subgeometric hypocoercivity for piecewise-deterministic Markov process Monte Carlo methods. *Electron. J. Probab.*, **26**, Paper No. 78, 26.
- Andrieu, Christophe, Lee, Anthony, Power, Sam, and Wang, Andi Q. 2022. Comparison of Markov chains via weak Poincaré inequalities with application to pseudo-marginal MCMC. *Ann. Statist.*, **50**(6), 3592–3618.
- Andrieu, Christophe, Lee, Anthony, Power, Sam, and Wang, Andi Q. 2024. Explicit convergence bounds for Metropolis Markov chains: isoperimetry, spectral gaps and profiles. *Ann. Appl. Probab.*
- Andrieu, Christophe, Lee, Anthony, Power, Sam, and Wang, Andi Q. 2026. Weak Poincaré inequalities for Markov chains: theory and applications. *Ann. Appl. Probab.*, **36**(1), 46–107.
- Arnaudon, Marc, Thalmaier, Anton, and Wang, Feng-Yu. 2006. Harnack inequality and heat kernel estimates on manifolds with curvature unbounded below. *Bull. Sci. Math.*, **130**(3), 223–233.
- Arnaudon, Marc, Bonnefont, Michel, and Joulin, Aldéric. 2018. Intertwinings and generalized Brascamp–Lieb inequalities. *Rev. Mat. Iberoam.*, **34**(3), 1021–1054.
- Ascolani, Filippo, Lavenant, Hugo, and Zanella, Giacomo. 2024. Entropy contraction of the Gibbs sampler under log-concavity. *arXiv preprint 2410.00858*.
- Baker, Jack, Fearnhead, Paul, Fox, Emily B., and Nemeth, Christopher. 2019. Control variates for stochastic gradient MCMC. *Stat. Comput.*, **29**(3), 599–615.
- Bakry, Dominique, Gentil, Ivan, and Ledoux, Michel. 2014. *Analysis and geometry of Markov diffusion operators*. Grundlehren der Mathematischen Wissenschaften [Fundamental Principles of Mathematical Sciences], vol. 348. Springer, Cham.
- Balasubramanian, Krishna, Chewi, Sinho, Erdogdu, Murat A., Salim, Adil, and Zhang, Matthew. 2022. Towards a theory of non-log-concave sampling: first-order stationarity guarantees for Langevin Monte Carlo. Pages 2896–2923 of: Loh, Po-Ling, and Raginsky, Maxim (eds), *Proceedings of Thirty Fifth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 178. PMLR.
- Bardet, Jean-Baptiste, Gozlan, Nathaël, Malrieu, Florent, and Zitt, Pierre-André. 2018. Functional inequalities for Gaussian convolutions of compactly supported measures: explicit bounds and dimension dependence. *Bernoulli*, **24**(1), 333–353.
- Barkhagen, Mathias, Chau, Huy N., Moulines, Éric, Rásonyi, Miklós, Sabanis, Sotirios, and Zhang, Ying. 2021. On stochastic gradient Langevin dynamics with dependent data streams in the logconcave case. *Bernoulli*, **27**(1), 1–33.
- Barthe, Franck, and Cordero-Erausquin, Dario. 2013. Invariances in variance estimates. *Proc. Lond. Math. Soc.* (3), **106**(1), 33–64.
- Baudoin, Fabrice. 2017. Bakry–Émery meet Villani. *Journal of Functional Analysis*, **273**(7), 2275–2291.
- Bauerschmidt, Roland, Bodineau, Thierry, and Dagallier, Benoit. 2024. Stochastic dynamics and the Polchinski equation: an introduction. *Probab. Surv.*, **21**, 200–290.
- Bauschke, Heinz H., Bolte, Jérôme, and Teboulle, Marc. 2017. A descent lemma beyond Lipschitz gradient continuity: first-order methods revisited and applications. *Math. Oper. Res.*, **42**(2), 330–348.
- Benamou, Jean-David, and Brenier, Yann. 2000. A computational fluid mechanics solution to the Monge–Kantorovich mass transfer problem. *Numer. Math.*, **84**(3), 375–393.
- Benton, Joe, Bortoli, Valentin De, Doucet, Arnaud, and Deligiannidis, George. 2024. Nearly d -linear convergence bounds for diffusion models via stochastic localization. In: *The Twelfth International Conference on Learning Representations*.
- Besag, Julian, Green, Peter, Higdon, David, and Mengersen, Kerrie. 1995. Bayesian computation and stochastic systems. *Statistical Science*, **10**(1), 3–66.
- Beskos, Alexandros, Pillai, Natesh, Roberts, Gareth, Sanz-Serna, Jesus-Maria, and Stuart, Andrew. 2013. Optimal tuning of the hybrid Monte Carlo algorithm. *Bernoulli*, **19**(5A), 1501–1534.

- Billera, Louis J., and Diaconis, Persi. 2001. A geometric interpretation of the Metropolis–Hastings algorithm. *Statist. Sci.*, **16**(4), 335–339.
- Blanchet, Adrien, and Bolte, Jérôme. 2018. A family of functional inequalities: Łojasiewicz inequalities and displacement convex functions. *J. Funct. Anal.*, **275**(7), 1650–1673.
- Block, Adam, Mroueh, Youssef, and Rakhlin, Alexander. 2022. Generative modeling with denoising auto-encoders and Langevin sampling. *arXiv preprint 2002.00107*.
- Bobkov, Sergey G. 1997. An isoperimetric inequality on the discrete cube, and an elementary proof of the isoperimetric inequality in Gauss space. *Ann. Probab.*, **25**(1), 206–214.
- Bobkov, Sergey G. 1999. Remarks on the Gromov–Milman inequality. *Vestn. Syktyvkar. Univ. Ser. I Mat. Mekh. Inform.*, 15–22.
- Bobkov, Sergey G., and Ledoux, Michel. 2000. From Brunn–Minkowski to Brascamp–Lieb and to logarithmic Sobolev inequalities. *Geom. Funct. Anal.*, **10**(5), 1028–1052.
- Bobkov, Sergey G., Gentil, Ivan, and Ledoux, Michel. 2001. Hypercontractivity of Hamilton–Jacobi equations. *J. Math. Pures Appl.* (9), **80**(7), 669–696.
- Bobkov, Serguei G., and Houdré, Christian. 1997. Some connections between isoperimetric and Sobolev-type inequalities. *Mem. Amer. Math. Soc.*, **129**(616), viii+111.
- Bolley, François, Gentil, Ivan, and Guillin, Arnaud. 2018. Dimensional improvements of the logarithmic Sobolev, Talagrand and Brascamp–Lieb inequalities. *Ann. Probab.*, **46**(1), 261–301.
- Borell, Christer. 1975. The Brunn–Minkowski inequality in Gauss space. *Invent. Math.*, **30**(2), 207–216.
- Borell, Christer. 1985. Geometric bounds on the Ornstein–Uhlenbeck velocity process. *Z. Wahrsch. Verw. Gebiete*, **70**(1), 1–13.
- Borell, Christer. 2000. Diffusion equations and geometric inequalities. *Potential Anal.*, **12**(1), 49–71.
- Bou-Rabee, Nawaf, Eberle, Andreas, and Zimmer, Raphael. 2020. Coupling and convergence for Hamiltonian Monte Carlo. *Ann. Appl. Probab.*, **30**(3), 1209–1250.
- Bou-Rabee, Nawaf, Mitra, Siddharth, and Wibisono, Andre. 2026. Tail-sensitive KL and Rényi convergence of unadjusted Hamiltonian Monte Carlo via one-shot couplings. *arXiv preprint 2601.09019*.
- Boucheron, Stéphane, Lugosi, Gábor, and Massart, Pascal. 2013. *Concentration inequalities*. Oxford University Press, Oxford. A nonasymptotic theory of independence, With a foreword by Michel Ledoux.
- Braverman, Mark, Hazan, Elad, Simchowitz, Max, and Woodworth, Blake. 2020. The gradient complexity of linear regression. Pages 627–647 of: Abernethy, Jacob, and Agarwal, Shivani (eds), *Proceedings of Thirty Third Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 125. PMLR.
- Brenier, Yann. 1991. Polar factorization and monotone rearrangement of vector-valued functions. *Comm. Pure Appl. Math.*, **44**(4), 375–417.
- Brigati, Giovanni. 2023. Time averages for kinetic Fokker–Planck equations. *Kinet. Relat. Models*, **16**(4), 524–539.
- Brigati, Giovanni, and Pedrotti, Francesco. 2024. Heat flow, log-concavity, and Lipschitz transport maps. *arXiv preprint 2404.15205*.
- Brigati, Giovanni, Stoltz, Gabriel, Wang, Andi Q., and Wang, Lihan. 2025. Explicit convergence rates of underdamped Langevin dynamics under weighted and weak Poincaré–Lions inequalities. *arXiv preprint 2407.16033*.
- Brunick, Gerard, and Shreve, Steven. 2013. Mimicking an Itô process by a solution of a stochastic differential equation. *Ann. Appl. Probab.*, **23**(4), 1584–1628.
- Bubeck, Sébastien. 2015. Convex optimization: algorithms and complexity. *Foundations and Trends® in Machine Learning*, **8**(3–4), 231–357.
- Bunne, Charlotte, Hsieh, Ya-Ping, Cuturi, Marco, and Krause, Andreas. 2023. The Schrödinger bridge between Gaussian measures has a closed form. Pages 5802–5833 of: Ruiz, Francisco, Dy, Jennifer, and van de Meent, Jan-Willem (eds), *Proceedings of the 26th International Conference on Artificial Intelligence and Statistics*. Proceedings of Machine Learning Research, vol. 206. PMLR.
- Burkholder, Donald L. 1973. Distribution function inequalities for martingales. *Ann. Probability*, **1**, 19–42.
- Caffarelli, Luis A. 2000. Monotonicity properties of optimal transportation and the FKG and related inequalities. *Comm. Math. Phys.*, **214**(3), 547–563.
- Cameron, Robert H., and Martin, William T. 1944. Transformations of Wiener integrals under translations. *Ann. of Math.* (2), **45**, 386–396.
- Cao, Yu, Lu, Jianfeng, and Lu, Yulong. 2019. Exponential decay of Rényi divergence under Fokker–Planck equations. *J. Stat. Phys.*, **176**(5), 1172–1184.

- Cao, Yu, Lu, Jianfeng, and Wang, Lihan. 2021. Complexity of randomized algorithms for underdamped Langevin dynamics. *Commun. Math. Sci.*, **19**(7), 1827–1853.
- Cao, Yu, Lu, Jianfeng, and Wang, Lihan. 2023. On explicit L^2 -convergence rate estimate for underdamped Langevin dynamics. *Arch. Ration. Mech. Anal.*, **247**(5), 90.
- Carmon, Yair, Duchi, John C., Hinder, Oliver, and Sidford, Aaron. 2020. Lower bounds for finding stationary points I. *Math. Program.*, **184**(1–2, Ser. A), 71–120.
- Carrillo, José A., and Vaes, Urbain. 2021. Wasserstein stability estimates for covariance-preconditioned Fokker–Planck equations. *Nonlinearity*, **34**(4), 2275–2295.
- Cattiaux, Patrick, Conforti, Giovanni, Gentil, Ivan, and Léonard, Christian. 2023. Time reversal of diffusion processes under a finite entropy condition. *Ann. Inst. Henri Poincaré Probab. Stat.*, **59**(4), 1844–1881.
- Chafaï, Djalil. 2004. Entropies, convexity, and functional inequalities: on Φ -entropies and Φ -Sobolev inequalities. *J. Math. Kyoto Univ.*, **44**(2), 325–363.
- Chafaï, Djalil, and Malrieu, Florent. 2010. On fine properties of mixtures with respect to concentration of measure and Sobolev type inequalities. *Ann. Inst. Henri Poincaré Probab. Stat.*, **46**(1), 72–96.
- Chatterji, Niladri, Flammarion, Nicolas, Ma, Yian, Bartlett, Peter, and Jordan, Michael. 2018. On the theory of variance reduction for stochastic gradient Monte Carlo. Pages 764–773 of: Dy, Jennifer, and Krause, Andreas (eds), *Proceedings of the 35th International Conference on Machine Learning*. Proceedings of Machine Learning Research, vol. 80. PMLR.
- Chatterji, Niladri S., Bartlett, Peter L., and Long, Philip M. 2022. Oracle lower bounds for stochastic gradient sampling algorithms. *Bernoulli*, **28**(2), 1074–1092.
- Chen, August Y., and Sridharan, Karthik. 2025. Optimization, isoperimetric inequalities, and sampling via Lyapunov potentials. *arXiv preprint 2410.02979*.
- Chen, Fan, Chewi, Sinho, Daskalakis, Constantinos, and Rakhlin, Alexander. 2026a. High-accuracy log-concave sampling with stochastic queries. *arXiv preprint 2602.14342*.
- Chen, Fan, Chewi, Sinho, Daskalakis, Constantinos, and Rakhlin, Alexander. 2026b. High-accuracy sampling for diffusion models and log-concave distributions. *arXiv preprint 2602.01338*.
- Chen, Gong, and Teboulle, Marc. 1993. Convergence analysis of a proximal-like minimization algorithm using Bregman functions. *SIAM J. Optim.*, **3**(3), 538–543.
- Chen, Hong-Bin, Chewi, Sinho, and Niles-Weed, Jonathan. 2021a. Dimension-free log-Sobolev inequalities for mixture distributions. *Journal of Functional Analysis*, **281**(11), 109236.
- Chen, Hongrui, Lee, Holden, and Lu, Jianfeng. 2023a. Improved analysis of score-based generative modeling: user-friendly bounds under minimal smoothness assumptions. Pages 4735–4763 of: Krause, Andreas, Brunskill, Emma, Cho, Kyunghyun, Engelhardt, Barbara, Sabato, Sivan, and Scarlett, Jonathan (eds), *Proceedings of the 40th International Conference on Machine Learning*. Proceedings of Machine Learning Research, vol. 202. PMLR.
- Chen, Sitan, Chewi, Sinho, Lee, Holden, Li, Yuanzhi, Lu, Jianfeng, and Salim, Adil. 2023b. The probability flow ODE is provably fast. Pages 68552–68575 of: Oh, A., Neumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds), *Advances in Neural Information Processing Systems*, vol. 36. Curran Associates, Inc.
- Chen, Sitan, Chewi, Sinho, Li, Jerry, Li, Yuanzhi, Salim, Adil, and Zhang, Anru R. 2023c. Sampling is as easy as learning the score: theory for diffusion models with minimal data assumptions. In: *The Eleventh International Conference on Learning Representations*.
- Chen, Tianrong, Liu, Guan-Horng, and Theodorou, Evangelos. 2022a. Likelihood training of Schrödinger bridge using forward-backward SDEs theory. In: *International Conference on Learning Representations*.
- Chen, Yongxin, Georgiou, Tryphon T., and Pavon, Michele. 2021b. Stochastic control liaisons: Richard Sinkhorn meets Gaspard Monge on a Schrödinger bridge. *SIAM Rev.*, **63**(2), 249–313.
- Chen, Yongxin, Chewi, Sinho, Salim, Adil, and Wibisono, Andre. 2022b. Improved analysis for a proximal algorithm for sampling. Pages 2984–3014 of: Loh, Po-Ling, and Raginsky, Maxim (eds), *Proceedings of Thirty Fifth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 178. PMLR.
- Chen, Yuansi. 2021. An almost constant lower bound of the isoperimetric coefficient in the KLS conjecture. *Geom. Funct. Anal.*, **31**(1), 34–61.
- Chen, Yuansi, and Eldan, Ronen. 2022. Localization schemes: a framework for proving mixing bounds for Markov chains (extended abstract). Pages 110–122 of: *2022 IEEE 63rd Annual Symposium on Foundations of Computer Science (FOCS)*.

- Chen, Yuansi, and Gtmiry, Khashayar. 2023a. A simple proof of the mixing of Metropolis-adjusted Langevin algorithm under smoothness and isoperimetry. *arXiv preprint 2304.04095*.
- Chen, Yuansi, and Gtmiry, Khashayar. 2023b. When does Metropolized Hamiltonian Monte Carlo provably outperform Metropolis-adjusted Langevin algorithm? *arXiv preprint 2304.04724*.
- Chen, Yuansi, Dwivedi, Raaz, Wainwright, Martin J., and Yu, Bin. 2020. Fast mixing of Metropolized Hamiltonian Monte Carlo: benefits of multi-step gradients. *J. Mach. Learn. Res.*, **21**, Paper No. 92, 71.
- Chen, Zongchen, and Vempala, Santosh S. 2019. Optimal convergence rate of Hamiltonian Monte Carlo for strongly logconcave distributions. Pages Art. No. 64, 12 of: *Approximation, randomization, and combinatorial optimization. Algorithms and techniques*. LIPIcs. Leibniz Int. Proc. Inform., vol. 145. Schloss Dagstuhl. Leibniz-Zent. Inform., Wadern.
- Cheng, Xiang, Chatterji, Niladri S., Bartlett, Peter L., and Jordan, Michael I. 2018. Underdamped Langevin MCMC: a non-asymptotic analysis. Pages 300–323 of: Bubeck, Sébastien, Perchet, Vianney, and Rigollet, Philippe (eds), *Proceedings of the 31st Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 75. PMLR.
- Cheng, Xiang, Wang, Bohan, Zhang, Jingzhao, and Zhu, Yusong. 2023. Fast conditional mixing of MCMC algorithms for non-log-concave distributions. Pages 13374–13394 of: Oh, A., Naumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds), *Advances in Neural Information Processing Systems*, vol. 36. Curran Associates, Inc.
- Chewi, Sinho. 2021. The entropic barrier is n -self-concordant. *arXiv e-prints*.
- Chewi, Sinho. 2023. *An optimization perspective on log-concave sampling and beyond*. Ph.D. thesis, MIT.
- Chewi, Sinho. 2024. *Log-concave sampling*. Forthcoming.
- Chewi, Sinho. 2025. *Lectures on optimization*. Available at https://chewisinho.github.io/opt_notes_final.pdf.
- Chewi, Sinho, and Pooladian, Aram-Alexandre. 2023. An entropic generalization of Caffarelli’s contraction theorem via covariance inequalities. *Reports. Mathematical*, **361**, 1471–1482.
- Chewi, Sinho, and Stromme, Austin J. 2024. The ballistic limit of the log-Sobolev constant equals the Polyak–Lojasiewicz constant. *arXiv preprint 2411.11415*.
- Chewi, Sinho, and Wibisono, Andre. 2026. Complexity of non-log-concave sampling in Fisher information. *arXiv preprint 2605.15859*.
- Chewi, Sinho, Le Gouic, Thibaut, Lu, Chen, Maunu, Tyler, Rigollet, Philippe, and Stromme, Austin J. 2020. Exponential ergodicity of mirror-Langevin diffusions. Pages 19573–19585 of: Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M. F., and Lin, H. (eds), *Advances in Neural Information Processing Systems*, vol. 33. Curran Associates, Inc.
- Chewi, Sinho, Lu, Chen, Ahn, Kwangjun, Cheng, Xiang, Gouic, Thibaut Le, and Rigollet, Philippe. 2021. Optimal dimension dependence of the Metropolis-adjusted Langevin algorithm. Pages 1260–1300 of: Belkin, Mikhail, and Kpotufe, Samory (eds), *Proceedings of Thirty Fourth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 134. PMLR.
- Chewi, Sinho, Gerber, Patrik R., Lu, Chen, Gouic, Thibaut Le, and Rigollet, Philippe. 2022. The query complexity of sampling from strongly log-concave distributions in one dimension. Pages 2041–2059 of: Loh, Po-Ling, and Raginsky, Maxim (eds), *Proceedings of Thirty Fifth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 178. PMLR.
- Chewi, Sinho, Gerber, Patrik R., Lee, Holden, and Lu, Chen. 2023a. Fisher information lower bounds for sampling. Pages 375–410 of: Agrawal, Shipra, and Orabona, Francesco (eds), *Proceedings of the 34th International Conference on Algorithmic Learning Theory*. Proceedings of Machine Learning Research, vol. 201. PMLR.
- Chewi, Sinho, De Dios Pont, Jaume, Li, Jerry, Lu, Chen, and Narayanan, Shyam. 2023b. Query lower bounds for log-concave sampling. Pages 2139–2148 of: *2023 IEEE 64th Annual Symposium on Foundations of Computer Science (FOCS)*.
- Chewi, Sinho, Nitanda, Atsushi, and Zhang, Matthew S. 2024. Uniform-in- N log-Sobolev inequality for the mean-field Langevin dynamics with convex energy. *arXiv preprint 2409.10440*.
- Chewi, Sinho, Erdogdu, Murat A., Li, Mufan, Shen, Ruoqi, and Zhang, Matthew S. 2025a. Analysis of Langevin Monte Carlo from Poincaré to log-Sobolev. *Found. Comput. Math.*, **25**(4), 1345–1395.
- Chewi, Sinho, Niles-Weed, Jonathan, and Rigollet, Philippe. 2025b. *Statistical optimal transport*. Lecture Notes in Mathematics, vol. 2364. Springer, Cham. École d’Été de Probabilités de Saint-Flour XLIX – 2019.
- Conforti, Giovanni, Durmus, Alain, and Gentiloni Silveri, Marta. 2025. KL convergence guarantees for score diffusion models under minimal data assumptions. *SIAM J. Math. Data Sci.*, **7**(1), 86–109.

- Cordero-Erausquin, Dario. 2002. Some applications of mass transport to Gaussian-type inequalities. *Arch. Ration. Mech. Anal.*, **161**(3), 257–269.
- Cordero-Erausquin, Dario. 2017. Transport inequalities for log-concave measures, quantitative forms, and applications. *Canad. J. Math.*, **69**(3), 481–501.
- Cordero-Erausquin, Dario, McCann, Robert J., and Schmuckenschläger, Michael. 2001. A Riemannian interpolation inequality à la Borell, Brascamp and Lieb. *Invent. Math.*, **146**(2), 219–257.
- Cordero-Erausquin, Dario, Fradelizi, Matthieu, and Maurey, Bernard. 2004. The (B) conjecture for the Gaussian measure of dilates of symmetric convex sets and related problems. *J. Funct. Anal.*, **214**(2), 410–427.
- Cousins, Ben, and Vempala, Santosh. 2018. Gaussian cooling and $O^*(n^3)$ algorithms for volume and Gaussian volume. *SIAM J. Comput.*, **47**(3), 1237–1273.
- Cover, Thomas M., and Thomas, Joy A. 2006. *Elements of information theory*. Second edn. Wiley-Interscience [John Wiley & Sons], Hoboken, NJ.
- Csiszár, Imre. 1967. Information-type measures of difference of probability distributions and indirect observations. *Studia Sci. Math. Hungar.*, **2**, 299–318.
- Cuesta, Juan Antonio, and Matrán, Carlos. 1989. Notes on the Wasserstein metric in Hilbert spaces. *Ann. Probab.*, **17**(3), 1264–1276.
- Cuturi, Marco. 2013. Sinkhorn distances: lightspeed computation of optimal transport. In: Burges, C.J., Bottou, L., Welling, M., Ghahramani, Z., and Weinberger, K.Q. (eds), *Advances in Neural Information Processing Systems*, vol. 26. Curran Associates, Inc.
- Daaloul, Chiheb, Le Gouic, Thibaut, Liandrat, Jacques, and Tournus, Magali. 2025. Convergence of the mirror Langevin algorithm on bounded domains under log-Sobolev inequality for the target measure. *hal-05127974*, 6.
- Dalalyan, Arnak S. 2017a. Further and stronger analogy between sampling and optimization: Langevin Monte Carlo and gradient descent. Pages 678–689 of: Kale, Satyen, and Shamir, Ohad (eds), *Proceedings of the 2017 Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 65. PMLR.
- Dalalyan, Arnak S. 2017b. Theoretical guarantees for approximate sampling from smooth and log-concave densities. *J. R. Stat. Soc. Ser. B. Stat. Methodol.*, **79**(3), 651–676.
- Dalalyan, Arnak S., and Karagulyan, Avetik. 2019. User-friendly guarantees for the Langevin Monte Carlo with inaccurate gradient. *Stochastic Process. Appl.*, **129**(12), 5278–5311.
- Dalalyan, Arnak S., and Karagulyan, Avetik. 2026. Improved guarantees for Langevin Monte Carlo with average smoothness. *arXiv preprint 2605.31413*.
- Dalalyan, Arnak S., and Riou-Durand, Lionel. 2020. On sampling from a log-concave density using kinetic Langevin diffusions. *Bernoulli*, **26**(3), 1956–1988.
- Dalalyan, Arnak S., and Tsybakov, Alexandre B. 2012. Sparse regression learning by aggregation and Langevin Monte-Carlo. *J. Comput. System Sci.*, **78**(5), 1423–1443.
- Davies, Ewan, Lee, Holden, Sandhu, Juspreet Singh, and Shi, Jonathan. 2026. Weak Poincaré inequalities via approximate stochastic localization: application to sampling the Sherrington–Kirkpatrick model. *arXiv preprint 2607.08160*.
- Davis, Burgess. 1976. On the L^p norms of stochastic integrals and other martingales. *Duke Math. J.*, **43**(4), 697–704.
- De Bortoli, Valentin. 2022. Convergence of denoising diffusion models under the manifold hypothesis. *Transactions on Machine Learning Research*.
- De Bortoli, Valentin, Thornton, James, Heng, Jeremy, and Doucet, Arnaud. 2021. Diffusion Schrödinger bridge with applications to score-based generative modeling. Pages 17695–17709 of: Ranzato, M., Beygelzimer, A., Dauphin, Y., Liang, P.S., and Vaughan, J. Wortman (eds), *Advances in Neural Information Processing Systems*, vol. 34. Curran Associates, Inc.
- Dembo, Amir, and Zeitouni, Ofer. 2010. *Large deviations techniques and applications*. Stochastic Modelling and Applied Probability, vol. 38. Springer-Verlag, Berlin. Corrected reprint of the second (1998) edition.
- Ding, Zhiyan, and Li, Qin. 2021. Langevin Monte Carlo: random coordinate descent and variance reduction. *J. Mach. Learn. Res.*, **22**, Paper No. 205, 51.
- Ding, Zhiyan, Li, Qin, Lu, Jianfeng, and Wright, Stephen J. 2021a. Random coordinate Langevin Monte Carlo. Pages 1683–1710 of: Belkin, Mikhail, and Kpotufe, Samory (eds), *Proceedings of Thirty Fourth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 134. PMLR.

- Ding, Zhiyan, Li, Qin, Lu, Jianfeng, and Wright, Stephen. 2021b. Random coordinate underdamped Langevin Monte Carlo. Pages 2701–2709 of: Banerjee, Arindam, and Fukumizu, Kenji (eds), *Proceedings of The 24th International Conference on Artificial Intelligence and Statistics*. Proceedings of Machine Learning Research, vol. 130. PMLR.
- do Carmo, Manfredo P. 1992. *Riemannian geometry*. Mathematics: Theory & Applications. Birkhäuser Boston, Inc., Boston, MA. Translated from the second Portuguese edition by Francis Flaherty.
- Dobson, Paul, and Bierkens, Joris. 2023. Infinite dimensional piecewise deterministic Markov processes. *Stochastic Process. Appl.*, **165**, 337–396.
- Dolbeault, Jean, Mouhot, Clément, and Schmeiser, Christian. 2009. Hypocoercivity for kinetic equations with linear relaxation terms. *Comptes Rendus Mathématique*, **347**(9-10), 511–516.
- Dolbeault, Jean, Mouhot, Clément, and Schmeiser, Christian. 2015. Hypocoercivity for linear kinetic equations conserving mass. *Trans. Amer. Math. Soc.*, **367**(6), 3807–3828.
- Duane, Simon, Kennedy, Anthony D., Pendleton, Brian J., and Roweth, Duncan. 1987. Hybrid Monte Carlo. *Phys. Lett. B*, **195**(2), 216–222.
- Dubey, Avinava, Reddi, Sashank J., Williamson, Sinead A., Póczos, Barnabás, Smola, Alexander J., and Xing, Eric P. 2016. Variance reduction in stochastic gradient Langevin dynamics. In: Lee, D., Sugiyama, M., Luxburg, U., Guyon, I., and Garnett, R. (eds), *Advances in Neural Information Processing Systems*, vol. 29. Curran Associates, Inc.
- Durmus, Alain, and Moulines, Éric. 2017. Nonasymptotic convergence analysis for the unadjusted Langevin algorithm. *Ann. Appl. Probab.*, **27**(3), 1551–1587.
- Durmus, Alain, and Moulines, Éric. 2019. High-dimensional Bayesian inference via the unadjusted Langevin algorithm. *Bernoulli*, **25**(4A), 2854–2882.
- Durmus, Alain, Majewski, Szymon, and Miasojedow, Błażej. 2019. Analysis of Langevin Monte Carlo via convex optimization. *J. Mach. Learn. Res.*, **20**, Paper No. 73, 46.
- Dwivedi, Raaz, Chen, Yuansi, Wainwright, Martin J., and Yu, Bin. 2019. Log-concave sampling: Metropolis–Hastings algorithms are fast. *Journal of Machine Learning Research*, **20**(183), 1–42.
- Dwork, Cynthia, and Roth, Aaron. 2013. The algorithmic foundations of differential privacy. *Found. Trends Theor. Comput. Sci.*, **9**(3-4), 211–487.
- Dyer, Martin, Frieze, Alan, and Kannan, Ravi. 1991. A random polynomial-time algorithm for approximating the volume of convex bodies. *J. Assoc. Comput. Mach.*, **38**(1), 1–17.
- Eberle, Andreas. 2011. Reflection coupling and Wasserstein contractivity without convexity. *C. R. Math. Acad. Sci. Paris*, **349**(19-20), 1101–1104.
- Eberle, Andreas. 2016. Reflection couplings and contraction rates for diffusions. *Probab. Theory Related Fields*, **166**(3-4), 851–886.
- Eberle, Andreas, and Lörler, Francis. 2024. Space-time divergence lemmas and optimal non-reversible lifts of diffusions on Riemannian manifolds with boundary. *arXiv preprint 2412.16710*.
- Eberle, Andreas, and Lörler, Francis. 2026. Non-reversible lifts of reversible diffusion processes and relaxation times. *Probab. Theory Related Fields*, **194**(1-2), 173–203.
- Eberle, Andreas, Guillin, Arnaud, and Zimmer, Raphael. 2019. Couplings and quantitative contraction rates for Langevin dynamics. *Ann. Probab.*, **47**(4), 1982–2010.
- Eberle, Andreas, Guillin, Arnaud, Hahn, Leo, Lörler, Francis, and Michel, Manon. 2025. Convergence of non-reversible Markov processes via lifting and flow Poincaré inequality. *arXiv preprint 2503.04238*.
- Eckhardt, Roger. 1987. Stan Ulam, John von Neumann, and the Monte Carlo method. With contributions by Tony Warnock, Gary D. Doolen and John Hendricks, Stanislaw Ulam 1909–1984.
- Edelman, Alan S. 1989. *Eigenvalues and condition numbers of random matrices*. Ph.D. thesis. Massachusetts Institute of Technology.
- Efron, Bradley, and Stein, Charles M. 1981. The jackknife estimate of variance. *Ann. Statist.*, **9**(3), 586–596.
- El Alaoui, Ahmed, Montanari, Andrea, and Sellke, Mark. 2022. Sampling from the Sherrington–Kirkpatrick Gibbs measure via algorithmic stochastic localization. Pages 323–334 of: *2022 IEEE 63rd Annual Symposium on Foundations of Computer Science—FOCS 2022*. IEEE Computer Soc., Los Alamitos, CA.
- Eldan, Ronen. 2013. Thin shell implies spectral gap up to polylog via a stochastic localization scheme. *Geom. Funct. Anal.*, **23**(2), 532–569.
- Eldan, Ronen. 2015. A two-sided estimate for the Gaussian noise stability deficit. *Invent. Math.*, **201**(2), 561–624.
- Eldan, Ronen. 2016. Skorokhod embeddings via stochastic flows on the space of Gaussian measures. *Ann. Inst. Henri Poincaré Probab. Stat.*, **52**(3), 1259–1280.

- Eldan, Ronen. 2018. Gaussian-width gradient complexity, reverse log-Sobolev inequalities and nonlinear large deviations. *Geom. Funct. Anal.*, **28**(6), 1548–1596.
- Eldan, Ronen. 2020. Taming correlations through entropy-efficient measure decompositions with applications to mean-field approximation. *Probab. Theory Related Fields*, **176**(3-4), 737–755.
- Eldan, Ronen. 2023. Analysis of high-dimensional distributions using pathwise methods. Pages 4246–4270 of: *ICM—International Congress of Mathematicians. Vol. 6. Sections 12–14*. EMS Press, Berlin.
- Eldan, Ronen, and Lee, James R. 2018. Regularization under diffusion and anticoncentration of the information content. *Duke Math. J.*, **167**(5), 969–993.
- Eldan, Ronen, and Shamir, Omer. 2022. Log concavity and concentration of Lipschitz functions on the Boolean hypercube. *J. Funct. Anal.*, **282**(8), Paper No. 109392, 22.
- Eldan, Ronen, Lee, James R., and Lehec, Joseph. 2017. Transport-entropy inequalities and curvature in discrete-space Markov chains. Pages 391–406 of: *A journey through discrete mathematics*. Springer, Cham.
- Eldan, Ronen, Mikulincer, Dan, and Zhai, Alex. 2020. The CLT in high dimensions: quantitative bounds via martingale embedding. *Ann. Probab.*, **48**(5), 2494–2524.
- Eldan, Ronen, Koehler, Frederic, and Zeitouni, Ofer. 2022. A spectral condition for spectral gap: fast mixing in high-temperature Ising models. *Probab. Theory Related Fields*, **182**(3-4), 1035–1051.
- Erbar, Matthias, and Maas, Jan. 2012. Ricci curvature of finite Markov chains via convexity of the entropy. *Arch. Ration. Mech. Anal.*, **206**(3), 997–1038.
- Erdogdu, Murat A., Mackey, Lester, and Shamir, Ohad. 2018. Global non-convex optimization with discretized diffusions. In: Bengio, S., Wallach, H., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds), *Advances in Neural Information Processing Systems*, vol. 31. Curran Associates, Inc.
- Erdogdu, Murat A., Hosseinzadeh, Rasa, and Zhang, Shunshi. 2022. Convergence of Langevin Monte Carlo in chi-squared and Rényi divergence. Pages 8151–8175 of: Camps-Valls, Gustau, Ruiz, Francisco J. R., and Valera, Isabel (eds), *Proceedings of The 25th International Conference on Artificial Intelligence and Statistics*. Proceedings of Machine Learning Research, vol. 151. PMLR.
- Evans, Lawrence C. 2010. *Partial differential equations*. Second edn. Graduate Studies in Mathematics, vol. 19. American Mathematical Society, Providence, RI.
- Eyring, Henry. 1935. The activated complex in chemical reactions. *J. Chem. Phys.*, 107–115.
- Fan, Jiaojiao, Yuan, Bo, and Chen, Yongxin. 2023. Improved dimension dependence of a proximal algorithm for sampling. Pages 1473–1521 of: Neu, Gergely, and Rosasco, Lorenzo (eds), *Proceedings of Thirty Sixth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 195. PMLR.
- Fan, Zexi, Li, Bowen, and Lu, Jianfeng. 2026. Sharp hypocoercive convergence estimates for underdamped Langevin dynamics via the modified L^2 method. *arXiv preprint 2604.10068*.
- Fathi, Max, and Shu, Yan. 2018. Curvature and transport inequalities for Markov chains in discrete spaces. *Bernoulli*, **24**(1), 672–698.
- Fathi, Max, Gozlan, Nathael, and Prod'homme, Maxime. 2020. A proof of the Caffarelli contraction theorem via entropic regularization. *Calc. Var. Partial Differential Equations*, **59**(3), Paper No. 96, 18.
- Fathi, Max, Mikulincer, Dan, and Shenfeld, Yair. 2024. Transportation onto log-Lipschitz perturbations. *Calc. Var. Partial Differential Equations*, **63**(3), Paper No. 61, 25.
- Folland, Gerald B. 1999. *Real analysis*. Second edn. Pure and Applied Mathematics (New York). John Wiley & Sons, Inc., New York. Modern techniques and their applications, A Wiley-Interscience Publication.
- Föllmer, Hans. 1985. An entropy approach to the time reversal of diffusion processes. Pages 156–163 of: *Stochastic differential systems (Marseille-Luminy, 1984)*. Lect. Notes Control Inf. Sci., vol. 69. Springer, Berlin.
- Foster, James, Lyons, Terry, and Oberhauser, Harald. 2021. The shifted ODE method for underdamped Langevin MCMC. *arXiv e-prints*.
- Ganesh, Arun, and Talwar, Kunal. 2020. Faster differentially private samplers via Rényi divergence analysis of discretized Langevin MCMC. Pages 7222–7233 of: Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M.F., and Lin, H. (eds), *Advances in Neural Information Processing Systems*, vol. 33. Curran Associates, Inc.
- Gangbo, Wilfrid, and McCann, Robert J. 1996. The geometry of optimal transportation. *Acta Math.*, **177**(2), 113–161.
- Gatmiry, Khashayar, and Vempala, Santosh S. 2022. Convergence of the Riemannian Langevin algorithm. *arXiv e-prints*.
- Gatmiry, Khashayar, Chen, Sitan, and Salim, Adil. 2026. High-accuracy and dimension-free sampling with diffusions. *arXiv preprint 2601.10708*.

- Ge, Rong, Lee, Holden, and Lu, Jianfeng. 2020. Estimating normalizing constants for log-concave distributions: algorithms and lower bounds. Pages 579–586 of: *STOC'20—Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing*. ACM, New York.
- Gentil, Ivan. 2008. From the Prékopa–Leindler inequality to modified logarithmic Sobolev inequality. *Ann. Fac. Sci. Toulouse Math.* (6), **17**(2), 291–308.
- Gentil, Ivan, Léonard, Christian, Ripani, Luigia, and Tamanini, Luca. 2020. An entropic interpolation proof of the HWI inequality. *Stochastic Process. Appl.*, **130**(2), 907–923.
- Girolami, Mark, and Calderhead, Ben. 2011. Riemann manifold Langevin and Hamiltonian Monte Carlo methods. *J. R. Stat. Soc. Ser. B Stat. Methodol.*, **73**(2), 123–214. With discussion and a reply by the authors.
- Gong, Yun, He, Niao, and Shen, Zebang. 2025. Poincare inequality for local log-Polyak–Łojasiewicz measures: non-asymptotic analysis in low-temperature regime. *arXiv preprint 2501.00429*.
- Gopi, Sivakanth, Lee, Yin Tat, and Liu, Daogao. 2022. Private convex optimization via exponential mechanism. Pages 1948–1989 of: Loh, Po-Ling, and Raginsky, Maxim (eds), *Proceedings of Thirty Fifth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 178. PMLR.
- Gopi, Sivakanth, Lee, Yin Tat, Liu, Daogao, Shen, Ruoyi, and Tian, Kevin. 2023. Algorithmic aspects of the log-Laplace transform and a non-Euclidean proximal sampler. Pages 2399–2439 of: Neu, Gergely, and Rosasco, Lorenzo (eds), *Proceedings of Thirty Sixth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 195. PMLR.
- Goyal, Alexander, Deligiannidis, George, and Kantas, Nikolas. 2026. Mixing time bounds for the Gibbs sampler under isoperimetry. *arXiv preprint 2506.22258*.
- Gozlan, Nathael, and Sylvestre, Maxime. 2025. Global regularity estimates for optimal transport via entropic regularisation. *arXiv preprint 2501.11382*.
- Gozlan, Nathael, Roberto, Cyril, and Samson, Paul-Marie. 2011. A new characterization of Talagrand’s transport-entropy inequalities and applications. *Ann. Probab.*, **39**(3), 857–880.
- Gozlan, Nathael, Roberto, Cyril, Samson, Paul-Marie, and Tetali, Prasad. 2014. Displacement convexity of entropy and related inequalities on graphs. *Probab. Theory Related Fields*, **160**(1-2), 47–94.
- Gross, Leonard. 1967. Abstract Wiener spaces. Pages 31–42 of: *Proc. Fifth Berkeley Sympos. Math. Statist. and Probability (Berkeley, Calif., 1965/66), Vol. II: Contributions to Probability Theory, Part I*. Univ. California Press, Berkeley, CA.
- Guédon, Olivier, and Milman, Emanuel. 2011. Interpolating thin-shell and sharp large-deviation estimates for isotropic log-concave measures. *Geom. Funct. Anal.*, **21**(5), 1043–1068.
- Guillin, Arnaud, and Wang, Feng-Yu. 2012. Degenerate Fokker–Planck equations: Bismut formula, gradient estimate and Harnack inequality. *J. Differential Equations*, **253**(1), 20–40.
- Guo, Dongning, Shamaï, Shlomo, and Verdú, Sergio. 2005. Mutual information and minimum mean-square error in Gaussian channels. *IEEE Trans. Inform. Theory*, **51**(4), 1261–1282.
- Gupta, Shivam, Cai, Linda, and Chen, Sitan. 2025. Faster diffusion sampling with randomized midpoints: sequential and parallel. Pages 97663–97698 of: Yue, Y., Garg, A., Peng, N., Sha, F., and Yu, R. (eds), *International Conference on Learning Representations*, vol. 2025.
- Gyöngy, István. 1986. Mimicking the one-dimensional marginal distributions of processes having an Itô differential. *Probab. Theory Relat. Fields*, **71**(4), 501–516.
- Hastings, W. Keith. 1970. Monte Carlo sampling methods using Markov chains and their applications. *Biometrika*, **57**(1), 97–109.
- Hausmann, Ulrich G., and Pardoux, Étienne. 1986. Time reversal of diffusions. *Ann. Probab.*, **14**(4), 1188–1205.
- He, Ye, Balasubramanian, Krishnakumar, and Erdogdu, Murat A. 2020. On the ergodicity, bias and asymptotic normality of randomized midpoint sampling method. Pages 7366–7376 of: Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M. F., and Lin, H. (eds), *Advances in Neural Information Processing Systems*, vol. 33. Curran Associates, Inc.
- Helfffer, Bernard, and Sjöstrand, Johannes. 1994. On the correlation for Kac-like models in the convex case. *J. Statist. Phys.*, **74**(1-2), 349–409.
- Ho, Jonathan, Jain, Ajay, and Abbeel, Pieter. 2020. Denoising diffusion probabilistic models. Pages 6840–6851 of: Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M.F., and Lin, H. (eds), *Advances in Neural Information Processing Systems*, vol. 33. Curran Associates, Inc.

- Hoffman, Matthew D., and Gelman, Andrew. 2014. The no-U-turn sampler: adaptively setting path lengths in Hamiltonian Monte Carlo. *J. Mach. Learn. Res.*, **15**, 1593–1623.
- Holley, Richard, and Stroock, Daniel. 1987. Logarithmic Sobolev inequalities and stochastic Ising models. *J. Statist. Phys.*, **46**(5-6), 1159–1194.
- Hörmander, Lars. 2007. *Notions of convexity*. Modern Birkhäuser Classics. Birkhäuser Boston, Inc., Boston, MA. Reprint of the 1994 edition.
- Hsieh, Ya-Ping, Kavis, Ali, Rolland, Paul, and Cevher, Volkan. 2018. Mirrored Langevin dynamics. In: Bengio, S., Wallach, H., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds), *Advances in Neural Information Processing Systems*, vol. 31. Curran Associates, Inc.
- Hsu, Elton P. 2002. *Stochastic analysis on manifolds*. Graduate Studies in Mathematics, vol. 38. American Mathematical Society, Providence, RI.
- Hu, Zhengmian, Huang, Feihu, and Huang, Heng. 2021. Optimal underdamped Langevin MCMC method. Pages 19363–19374 of: Ranzato, M., Beygelzimer, A., Dauphin, Y., Liang, P.S., and Vaughan, J. Wortman (eds), *Advances in Neural Information Processing Systems*, vol. 34. Curran Associates, Inc.
- Huang, Zhihan, Wei, Yuting, and Chen, Yuxin. 2026. Denoising diffusion probabilistic models are optimally adaptive to unknown low dimensionality. *Mathematics of Operations Research*, **0**(0).
- Hyvärinen, Aapo. 2005. Estimation of non-normalized statistical models by score matching. *J. Mach. Learn. Res.*, **6**, 695–709.
- Isaksson, Marcus, and Mossel, Elchanan. 2012. Maximally stable Gaussian partitions with discrete applications. *Israel J. Math.*, **189**, 347–396.
- Jain, Nishant, and Zhang, Tong. 2026. A sharp KL convergence analysis for diffusion models under minimal assumptions. In: *The Fourteenth International Conference on Learning Representations*.
- Janati, Hicham, Muzellec, Boris, Peyré, Gabriel, and Cuturi, Marco. 2020. Entropic optimal transport between unbalanced Gaussian measures has a closed form. Pages 10468–10479 of: Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M.F., and Lin, H. (eds), *Advances in Neural Information Processing Systems*, vol. 33. Curran Associates, Inc.
- Jeon, Moongyu, Shin, Sangwoo, Jeon, Dongjae, and No, Albert. 2025. Information-theoretic discrete diffusion. Pages 18378–18404 of: Belgrave, D., Zhang, C., Lin, H., Pascanu, R., Koniusz, P., Ghassemi, M., and Chen, N. (eds), *Advances in Neural Information Processing Systems*, vol. 38. Curran Associates, Inc.
- Jiang, Qijia. 2021. Mirror Langevin Monte Carlo: the case under isoperimetry. Pages 715–725 of: Ranzato, M., Beygelzimer, A., Dauphin, Y., Liang, P.S., and Vaughan, J. Wortman (eds), *Advances in Neural Information Processing Systems*, vol. 34. Curran Associates, Inc.
- Johnson, Rie, and Zhang, Tong. 2013. Accelerating stochastic gradient descent using predictive variance reduction. In: Burges, C.J., Bottou, L., Welling, M., Ghahramani, Z., and Weinberger, K.Q. (eds), *Advances in Neural Information Processing Systems*, vol. 26. Curran Associates, Inc.
- Jordan, Richard, Kinderlehrer, David, and Otto, Felix. 1998. The variational formulation of the Fokker–Planck equation. *SIAM J. Math. Anal.*, **29**(1), 1–17.
- Kandasamy, Saravanan, and Nagaraj, Dheeraj. 2024. The Poisson midpoint method for Langevin dynamics: provably efficient discretization for diffusion models. Pages 65972–66024 of: Globerson, A., Mackey, L., Belgrave, D., Fan, A., Paquet, U., Tomczak, J., and Zhang, C. (eds), *Advances in Neural Information Processing Systems*, vol. 37. Curran Associates, Inc.
- Kannan, Ravi, Lovász, László, and Simonovits, Miklós. 1995. Isoperimetric problems for convex bodies and a localization lemma. *Discrete Comput. Geom.*, **13**(3-4), 541–559.
- Kannan, Ravi, Lovász, László, and Simonovits, Miklós. 1997. Random walks and an $O^*(n^5)$ volume algorithm for convex bodies. *Random Structures Algorithms*, **11**(1), 1–50.
- Kannan, Ravi, Lovász, László, and Montenegro, Ravi. 2006. Blocking conductance and mixing in random walks. *Combin. Probab. Comput.*, **15**(4), 541–570.
- Kantorovitch, Leonid V. 1942. On the translocation of masses. *C. R. (Doklady) Acad. Sci. URSS (N.S.)*, **37**, 199–201.
- Karimi, Hamed, Nutini, Julie, and Schmidt, Mark. 2016. Linear convergence of gradient and proximal-gradient methods under the Polyak–Lojasiewicz condition. Pages 795–811 of: *European Conference on Machine Learning and Knowledge Discovery in Databases—Volume 9851*. ECML PKDD 2016. Berlin, Heidelberg: Springer-Verlag.

- Kearns, Michael, Mansour, Yishay, Ron, Dana, Rubinfeld, Ronitt, Schapire, Robert E., and Sellie, Linda. 1994. On the learnability of discrete distributions. Pages 273–282 of: *Proceedings of the Twenty-Sixth Annual ACM Symposium on Theory of Computing*. STOC '94. New York, NY, USA: Association for Computing Machinery.
- Kim, Kyurae, Gruffaz, Samuel, Park, Ji Won, and Durmus, Alain O. 2025. Analysis of kinetic Langevin Monte Carlo under the stochastic exponential Euler discretization from underdamped all the way to overdamped. *arXiv preprint 2510.03949*.
- Kim, Young-Heon, and Milman, Emanuel. 2012. A generalization of Caffarelli's contraction theorem via (reverse) heat flow. *Math. Ann.*, **354**(3), 827–862.
- Kindler, Guy, Kirshner, Naomi, and O'Donnell, Ryan. 2018. Gaussian noise sensitivity and Fourier tails. *Israel J. Math.*, **225**(1), 71–109.
- Kinoshita, Yuri, and Suzuki, Taiji. 2022. Improved convergence rate of stochastic gradient Langevin dynamics with variance reduction and its application to optimization. Pages 19022–19034 of: Koyejo, S., Mohamed, S., Agarwal, A., Belgrave, D., Cho, K., and Oh, A. (eds), *Advances in Neural Information Processing Systems*, vol. 35. Curran Associates, Inc.
- Klartag, Bo'az. 2018. Eldan's stochastic localization and tubular neighborhoods of complex-analytic sets. *J. Geom. Anal.*, **28**(3), 2008–2027.
- Klartag, Bo'az. 2023. Logarithmic bounds for isoperimetry and slices of convex sets. *Ars Inven. Anal.*, Paper No. 4, 17.
- Klartag, Bo'az, and Lehec, Joseph. 2022. Bourgain's slicing problem and KLS isoperimetry up to polylog. *Geom. Funct. Anal.*, **32**(5), 1134–1159.
- Klartag, Boaz, and Lehec, Joseph. 2026. Thin-shell bounds via parallel coupling. *arXiv preprint 2507.15495*.
- Klartag, Bo'az, and Putterman, Eli. 2023. Spectral monotonicity under Gaussian convolution. *Ann. Fac. Sci. Toulouse Math.* (6), **32**(5), 939–967.
- Klartag, Bo'az, Kozma, Gady, Ralli, Peter, and Tetali, Prasad. 2016. Discrete curvature and abelian groups. *Canad. J. Math.*, **68**(3), 655–674.
- Kolesnikov, Alexander V. 2014. Hessian metrics, $CD(K, N)$ -spaces, and optimal transportation of log-concave measures. *Discrete Contin. Dyn. Syst.*, **34**(4), 1511–1532.
- Kong, Xianghao, Brekelmans, Rob, and Ver Steeg, Greg. 2023. Information-theoretic diffusion. In: *The Eleventh International Conference on Learning Representations*.
- Kook, Yunbum. 2025. Zeroth-order log-concave sampling. *arXiv preprint 2507.18021*.
- Kook, Yunbum, and Vempala, Santosh S. 2025a. Faster logconcave sampling from a cold start in high dimension. *arXiv preprint 2505.01937*.
- Kook, Yunbum, and Vempala, Santosh S. 2025b. The localization method for high-dimensional inequalities. *arXiv preprint 2512.10848*.
- Kook, Yunbum, and Vempala, Santosh S. 2025c. Sampling and integration of logconcave functions by algorithmic diffusion. Pages 924–932 of: *Proceedings of the 57th Annual ACM Symposium on Theory of Computing*. STOC '25. New York, NY, USA: Association for Computing Machinery.
- Kook, Yunbum, and Vempala, Santosh S. 2026. A unified complexity bound for logconcave sampling. *arXiv preprint 2606.12694*.
- Kook, Yunbum, and Zhang, Matthew S. 2024. Covariance estimation using Markov chain Monte Carlo. *arXiv preprint 2410.17147*.
- Kook, Yunbum, and Zhang, Matthew S. 2025. Rényi-infinity constrained sampling with d^3 membership queries. Pages 5278–5306.
- Kook, Yunbum, Vempala, Santosh S., and Zhang, Matthew S. 2024. In-and-out: algorithmic diffusion for sampling convex bodies. Pages 108354–108388 of: Globerson, A., Mackey, L., Belgrave, D., Fan, A., Paquet, U., Tomczak, J., and Zhang, C. (eds), *Advances in Neural Information Processing Systems*, vol. 37. Curran Associates, Inc.
- Kramers, Hendrik A. 1940. Brownian motion in a field of force and the diffusion model of chemical reactions. *Physica*, **7**, 284–304.
- Kuntz, Juan, Ottobre, Michela, and Stuart, Andrew M. 2018. Non-stationary phase of the MALA algorithm. *Stoch. Partial Differ. Equ. Anal. Comput.*, **6**(3), 446–499.
- Latala, Rafa l, and Oleszkiewicz, Krzysztof. 2000. Between Sobolev and Poincaré. Pages 147–168 of: *Geometric aspects of functional analysis*. Lecture Notes in Math., vol. 1745. Springer, Berlin.
- Lawler, Gregory F., and Sokal, Alan D. 1988. Bounds on the L^2 spectrum for Markov chains and Markov processes: a generalization of Cheeger's inequality. *Trans. Amer. Math. Soc.*, **309**(2), 557–580.

- Le Gall, Jean-François. 2016. *Brownian motion, martingales, and stochastic calculus*. French edn. Graduate Texts in Mathematics, vol. 274. Springer, [Cham].
- Ledoux, Michel. 2000. The geometry of Markov diffusion generators. vol. 9. Probability theory.
- Ledoux, Michel. 2001. *The concentration of measure phenomenon*. Mathematical Surveys and Monographs, vol. 89. American Mathematical Society, Providence, RI.
- Ledoux, Michel. 2018. *Remarks on some transportation cost inequalities*.
- Lee, Holden, Lu, Jianfeng, and Tan, Yixin. 2022. Convergence for score-based generative modeling with polynomial complexity. In: Oh, Alice H., Agarwal, Alekh, Belgrave, Danielle, and Cho, Kyunghyun (eds), *Advances in Neural Information Processing Systems*.
- Lee, Holden, Lu, Jianfeng, and Tan, Yixin. 2023. Convergence of score-based generative modeling for general data distributions. Pages 946–985 of: Agrawal, Shipra, and Orabona, Francesco (eds), *Proceedings of the 34th International Conference on Algorithmic Learning Theory*. Proceedings of Machine Learning Research, vol. 201. PMLR.
- Lee, Hyun S. 2025. *Scaling for MCMC on heterogeneous targets and nonstationary phases*. Ph.D. thesis, University of Warwick.
- Lee, Yin Tat, and Vempala, Santosh S. 2017. Eldan’s stochastic localization and the KLS hyperplane conjecture: an improved lower bound for expansion. Pages 998–1007 of: *58th Annual IEEE Symposium on Foundations of Computer Science—FOCS 2017*. IEEE Computer Soc., Los Alamitos, CA.
- Lee, Yin Tat, and Vempala, Santosh S. 2018. Stochastic localization + Stieltjes barrier = tight bound for log-Sobolev. Pages 1122–1129 of: *STOC’18—Proceedings of the 50th Annual ACM SIGACT Symposium on Theory of Computing*. ACM, New York.
- Lee, Yin Tat, and Yue, Man-Chung. 2021. Universal barrier is n -self-concordant. *Math. Oper. Res.*, **46**(3), 1129–1148.
- Lee, Yin Tat, Shen, Ruoqi, and Tian, Kevin. 2020. Logsmooth gradient concentration and tighter runtimes for Metropolized Hamiltonian Monte Carlo. Pages 2565–2597 of: Abernethy, Jacob, and Agarwal, Shivani (eds), *Proceedings of Thirty Third Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 125. PMLR.
- Lee, Yin Tat, Shen, Ruoqi, and Tian, Kevin. 2021a. Lower bounds on Metropolized sampling methods for well-conditioned distributions. Pages 18812–18824 of: Ranzato, M., Beygelzimer, A., Dauphin, Y., Liang, P.S., and Vaughan, J. Wortman (eds), *Advances in Neural Information Processing Systems*, vol. 34. Curran Associates, Inc.
- Lee, Yin Tat, Shen, Ruoqi, and Tian, Kevin. 2021b. Structured logconcave sampling with a restricted Gaussian oracle. Pages 2993–3050 of: Belkin, Mikhail, and Kpotufe, Samory (eds), *Proceedings of Thirty Fourth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 134. PMLR.
- Lee, Yin Tat, Shen, Ruoqi, and Tian, Kevin. 2021c. Structured logconcave sampling with a restricted Gaussian oracle. *arXiv e-prints*.
- Lehec, Joseph. 2013. Representation formula for the entropy and functional inequalities. *Ann. Inst. Henri Poincaré Probab. Stat.*, **49**(3), 885–899.
- Lehec, Joseph. 2025. Convergence in total variation for the kinetic Langevin algorithm. *Math. Stat. Learn.*, **8**(1-2), 71–104.
- Léonard, Christian. 2012. From the Schrödinger problem to the Monge–Kantorovich problem. *J. Funct. Anal.*, **262**(4), 1879–1920.
- Léonard, Christian. 2017. On the convexity of the entropy along entropic interpolations. Pages 194–242 of: *Measure theory in non-smooth spaces*. Partial Differ. Equ. Meas. Theory. De Gruyter Open, Warsaw.
- Li, Bowen, and Lu, Jianfeng. 2026. Space-time log-Sobolev inequality and hypocoercive hypercontractivity for underdamped Langevin dynamics. *arXiv preprint 2605.25083*.
- Li, Gen, and Cai, Changxiao. 2025. Provable acceleration for diffusion models under minimal assumptions. *arXiv preprint 2410.23285*.
- Li, Gen, and Jiao, Yuchen. 2025. Improved convergence rate for diffusion probabilistic models. In: *The Thirteenth International Conference on Learning Representations*.
- Li, Gen, and Yan, Yuling. 2025. $O(d/T)$ convergence theory for diffusion probabilistic models under minimal assumptions. In: *The Thirteenth International Conference on Learning Representations*.
- Li, Gen, Huang, Yu, Efimov, Timofey, Wei, Yuting, Chi, Yuejie, and Chen, Yuxin. 2024a. Accelerating convergence of score-based diffusion models, provably. Pages 27942–27954 of: Salakhutdinov, Ruslan, Kolter, Zico, Heller,

- Katherine, Weller, Adrian, Oliver, Nuria, Scarlett, Jonathan, and Berkenkamp, Felix (eds), *Proceedings of the 41st International Conference on Machine Learning*. Proceedings of Machine Learning Research, vol. 235. PMLR.
- Li, Gen, Wei, Yuting, Chi, Yuejie, and Chen, Yuxin. 2024b. A sharp convergence theory for the probability flow ODEs of diffusion models. *arXiv preprint 2408.02320*.
- Li, Lei, Wang, Mengchao, and Wang, Yuliang. 2025a. Estimates of the numerical density for stochastic differential equations with multiplicative noise. *arXiv preprint 2409.04991*.
- Li, Mufan (Bill), and Erdogdu, Murat A. 2023. Riemannian Langevin algorithm for solving semidefinite programs. *Bernoulli*, **29**(4), 3093–3113.
- Li, Ruilin, Tao, Molei, Vempala, Santosh S., and Wibisono, Andre. 2022a. The mirror Langevin algorithm converges with vanishing bias. Pages 718–742 of: Dasgupta, Sanjoy, and Haghtalab, Nika (eds), *Proceedings of the 33rd International Conference on Algorithmic Learning Theory*. Proceedings of Machine Learning Research, vol. 167. PMLR.
- Li, Ruilin, Zha, Hongyuan, and Tao, Molei. 2022b. $\sqrt{d}$ dimension dependence of Langevin Monte Carlo. In: *International Conference on Learning Representations*.
- Li, Runjia, Di, Qiwei, and Gu, Quanquan. 2025b. Unified convergence analysis for score-based diffusion models with deterministic samplers. Pages 15599–15655 of: Yue, Y., Garg, A., Peng, N., Sha, F., and Yu, R. (eds), *International Conference on Learning Representations*, vol. 2025.
- Li, Xuechen, Wu, Yi, Mackey, Lester, and Erdogdu, Murat A. 2019. Stochastic Runge–Kutta accelerates Langevin Monte Carlo and beyond. In: Wallach, H., Larochelle, H., Beygelzimer, A., d'Alché-Buc, F., Fox, E., and Garnett, R. (eds), *Advances in Neural Information Processing Systems*, vol. 32. Curran Associates, Inc.
- Liang, Jiadong, Huang, Zhihan, and Chen, Yuxin. 2025. Low-dimensional adaptation of diffusion models: convergence in total variation (extended abstract). Pages 3723–3729 of: Haghtalab, Nika, and Moitra, Ankur (eds), *Proceedings of Thirty Eighth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 291. PMLR.
- Liang, Jiaming, and Chen, Yongxin. 2022a. A proximal algorithm for sampling. *arXiv e-prints*.
- Liang, Jiaming, and Chen, Yongxin. 2022b. A proximal algorithm for sampling from non-smooth potentials. *arXiv e-prints*.
- Liang, Jiaming, Mitra, Siddharth, and Wibisono, Andre. 2024. On independent samples along the Langevin diffusion and the unadjusted Langevin algorithm. *arXiv preprint 2402.17067*.
- Lipman, Yaron, Chen, Ricky T. Q., Ben-Hamu, Heli, Nickel, Maximilian, and Le, Matthew. 2023. Flow matching for generative modeling. In: *The Eleventh International Conference on Learning Representations*.
- Liu, Linghai, and Chewi, Sinho. 2026. A proximal gradient algorithm for composite log-concave sampling. *arXiv preprint 2605.12461*.
- Liu, Xingchao, Wu, Lemeng, Ye, Mao, and Liu, Qiang. 2022. Let us build bridges: understanding and extending diffusion generative models. *arXiv preprint 2208.14699*.
- Liu, Xingchao, Gong, Chengyue, and Liu, Qiang. 2023. Flow straight and fast: learning to generate and transfer data with rectified flow. In: *The Eleventh International Conference on Learning Representations*.
- Liu, Yuhao, Chen, Yu, Hu, Rui, and Huang, Longbo. 2025. Finite-time analysis of discrete-time stochastic interpolants. In: *Forty-second International Conference on Machine Learning*.
- Lott, John, and Villani, Cédric. 2009. Ricci curvature for metric-measure spaces via optimal transport. *Ann. of Math.* (2), **169**(3), 903–991.
- Lovász, László, and Simonovits, Miklós. 1993. Random walks in a convex body and an improved volume algorithm. *Random Structures Algorithms*, **4**(4), 359–412.
- Lu, Haihao, Freund, Robert M., and Nesterov, Yurii. 2018. Relatively smooth convex optimization by first-order methods, and applications. *SIAM J. Optim.*, **28**(1), 333–354.
- Lu, Jianfeng. 2025. Quantitative hypocoercivity and lifting of classical and quantum dynamics. *arXiv preprint 2510.22305*.
- Lu, Jianfeng. 2026. A sharp hypocoercive entropy decay estimate for underdamped Langevin dynamics. *arXiv preprint 2605.01933*.
- Lu, Jianfeng, and Wang, Lihan. 2022a. Complexity of zigzag sampling algorithm for strongly log-concave distributions. *Stat. Comput.*, **32**(3), Paper No. 48, 12.
- Lu, Jianfeng, and Wang, Lihan. 2022b. On explicit L^2 -convergence rate estimate for piecewise deterministic Markov processes in MCMC algorithms. *Ann. Appl. Probab.*, **32**(2), 1333–1361.

- Lu, Jianfeng, and Wang, Yuliang. 2025. Long-time reverse transportation inequalities for non-globally-dissipative Langevin dynamics. *arXiv preprint 2512.18598*.
- Lu, Jianfeng, Ye, Xuda, and Zhou, Zhennan. 2025. Mean square error analysis of stochastic gradient and variance-reduced sampling algorithms. *arXiv preprint 2511.04413*.
- Ma, Yi-An, Chatterji, Niladri S., Cheng, Xiang, Flammarion, Nicolas, Bartlett, Peter L., and Jordan, Michael I. 2021. Is there an analog of Nesterov acceleration for gradient-based MCMC? *Bernoulli*, **27**(3), 1942 – 1992.
- Maas, Jan. 2011. Gradient flows of the entropy for finite Markov chains. *J. Funct. Anal.*, **261**(8), 2250–2292.
- Mallasto, Anton, Gerolin, Augusto, and Minh, Hà Quang. 2022. Entropy-regularized 2-Wasserstein distance between Gaussian measures. *Inf. Geom.*, **5**(1), 289–323.
- Mangoubi, Oren, and Smith, Aaron. 2019. Mixing of Hamiltonian Monte Carlo on strongly log-concave distributions 2: numerical integrators. Pages 586–595 of: Chaudhuri, Kamalika, and Sugiyama, Masashi (eds), *Proceedings of the Twenty-Second International Conference on Artificial Intelligence and Statistics*. Proceedings of Machine Learning Research, vol. 89. PMLR.
- Mangoubi, Oren, and Vishnoi, Nisheeth. 2018. Dimensionally tight bounds for second-order Hamiltonian Monte Carlo. In: Bengio, S., Wallach, H., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds), *Advances in Neural Information Processing Systems*, vol. 31. Curran Associates, Inc.
- Marton, Katalin. 1996. Bounding $\bar{d}$ -distance by informational divergence: a method to prove measure concentration. *Ann. Probab.*, **24**(2), 857–866.
- Maunu, Tyler, and Yao, Jiayi. 2025. Subspace Langevin Monte Carlo. *arXiv preprint 2412.13928*.
- McSherry, Frank, and Talwar, Kunal. 2007. Mechanism design via differential privacy. In: *48th Annual IEEE Symposium on Foundations of Computer Science (FOCS'07)*.
- Menz, Georg, and Schlichting, André. 2014. Poincaré and logarithmic Sobolev inequalities by decomposition of the energy landscape. *Ann. Probab.*, **42**(5), 1809–1884.
- Metropolis, Nicholas, Rosenbluth, Arianna W., Rosenbluth, Marshall N., Teller, Augusta H., and Teller, Edward. 1953. Equation of state calculations by fast computing machines. *The Journal of Chemical Physics*, **21**(6), 1087–1092.
- Mielke, Alexander. 2013. Geodesic convexity of the relative entropy in reversible Markov chains. *Calc. Var. Partial Differential Equations*, **48**(1-2), 1–31.
- Mikulincer, Dan, and Shenfeld, Yair. 2023. On the Lipschitz properties of transportation along heat flows. Pages 269–290 of: *Geometric aspects of functional analysis*. Lecture Notes in Math., vol. 2327. Springer, Cham.
- Mikulincer, Dan, and Shenfeld, Yair. 2024. The Brownian transport map. *Probab. Theory Related Fields*, **190**(1-2), 379–444.
- Milman, Emanuel. 2009. On the role of convexity in isoperimetry, spectral gap and concentration. *Invent. Math.*, **177**(1), 1–43.
- Milstein, Grigori N., and Tretyakov, Michael V. 2021. *Stochastic numerics for mathematical physics*. Second edn. Scientific Computation. Springer, Cham.
- Mironov, Ilya. 2017. Rényi differential privacy. Pages 263–275 of: *2017 IEEE 30th Computer Security Foundations Symposium (CSF)*.
- Mitra, Siddharth, Srinivasan, Vishwak, Wang, Xiuyuan, and Wibisono, Andre. 2026. Accelerated mixing time of randomized Hamiltonian Monte Carlo. *arXiv preprint 2607.12902*.
- Miyasawa, Koichi. 1961. An empirical Bayes estimator of the mean of a normal population. *Bull. Inst. Internat. Statist.*, **38**, 181–188.
- Monge, Gaspard. 1781. Mémoire sur la théorie des déblais et des remblais. *Histoire de l'Académie Royale des Sciences de Paris, avec les Mémoires de Mathématique et de Physique pour la même année*, 666–704.
- Monmarché, Pierre, and Wang, Lihan. 2026. On the entropic convergence for piecewise deterministic samplers: speedup and obstruction. *arXiv preprint 2606.26086*.
- Monmarché, Pierre. 2026. Nesterov acceleration for the Wasserstein minimization of displacement-convex free energies. *arXiv preprint 2605.13186*.
- Montanari, Andrea. 2023. Sampling, diffusions, and stochastic localization. *arXiv preprint 2305.10690*.
- Mossel, Elchanan, and Neeman, Joe. 2015. Robust optimality of Gaussian noise stability. *J. Eur. Math. Soc. (JEMS)*, **17**(2), 433–482.
- Mou, Wenlong, Flammarion, Nicolas, Wainwright, Martin J., and Bartlett, Peter L. 2022. Improved bounds for discretization of Langevin diffusions: near-optimal rates without convexity. *Bernoulli*, **28**(3), 1577–1601.

- Neal, Radford M. 2011. MCMC using Hamiltonian dynamics. Pages 113–162 of: *Handbook of Markov chain Monte Carlo*. Chapman & Hall/CRC Handb. Mod. Stat. Methods. CRC Press, Boca Raton, FL.
- Negrea, Jeffrey. 2022. *Approximations and scaling limits of Markov chains with applications to MCMC and approximate inference*. Ann Arbor, MI: ProQuest LLC. Thesis (Ph.D.)—University of Toronto (Canada).
- Nemirovsky, Arkadii S., and Yudin, David B. 1983. *Problem complexity and method efficiency in optimization*. Wiley-Interscience Series in Discrete Mathematics. John Wiley & Sons, Inc., New York. Translated from the Russian and with a preface by E. R. Dawson.
- Nesterov, Yu. E. 1983. A method for solving the convex programming problem with convergence rate $O(1/k^2)$. *Dokl. Akad. Nauk SSSR*, **269**(3), 543–547.
- Nesterov, Yurii. 2018. *Lectures on convex optimization*. Springer Optimization and Its Applications, vol. 137. Springer, Cham.
- Nesterov, Yurii, and Nemirovskii, Arkadii S. 1994. *Interior-point polynomial algorithms in convex programming*. SIAM Studies in Applied Mathematics, vol. 13. Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA.
- Nickl, Richard, and Wang, Sven. 2024. On polynomial-time computation of high-dimensional posterior measures by Langevin-type algorithms. *J. Eur. Math. Soc. (JEMS)*, **26**(3), 1031–1112.
- Niles-Weed, Jonathan. 2026. Reweighted information inequalities. *arXiv preprint 2603.13135*.
- Nishimura, Akihiko, and Suchard, Marc A. 2023. Prior-preconditioned conjugate gradient method for accelerated Gibbs sampling in “large n , large p ” Bayesian sparse regression. *J. Amer. Statist. Assoc.*, **118**(544), 2468–2481.
- Nutz, Marcel, and Wiesel, Johannes. 2022. Entropic optimal transport: convergence of potentials. *Probab. Theory Related Fields*, **184**(1-2), 401–424.
- Ollivier, Yann. 2007. Ricci curvature of metric spaces. *C. R. Math. Acad. Sci. Paris*, **345**(11), 643–646.
- Ollivier, Yann. 2009. Ricci curvature of Markov chains on metric spaces. *J. Funct. Anal.*, **256**(3), 810–864.
- Ollivier, Yann, and Villani, Cédric. 2012. A curved Brunn–Minkowski inequality on the discrete hypercube, or: what is the Ricci curvature of the discrete hypercube? *SIAM J. Discrete Math.*, **26**(3), 983–996.
- Otto, Felix. 2001. The geometry of dissipative evolution equations: the porous medium equation. *Comm. Partial Differential Equations*, **26**(1-2), 101–174.
- Otto, Felix, and Villani, Cédric. 2000. Generalization of an inequality by Talagrand and links with the logarithmic Sobolev inequality. *J. Funct. Anal.*, **173**(2), 361–400.
- Otto, Felix, and Villani, Cédric. 2001. Comment on: “Hypercontractivity of Hamilton–Jacobi equations” [J. Math. Pures Appl. (9) **80** (2001), no. 7, 669–696] by S. G. Bobkov, I. Gentil and M. Ledoux. *J. Math. Pures Appl.* (9), **80**(7), 697–700.
- Paulin, Daniel, and Whalley, Peter A. 2026. Theoretical guarantees for stochastic gradient sampling methods via Gaussian convolution inequalities. *arXiv preprint 2604.24632*.
- Pavliotis, Grigorios A. 2014. *Stochastic processes and applications*. Texts in Applied Mathematics, vol. 60. Springer, New York. Diffusion processes, the Fokker–Planck and Langevin equations.
- Pidstrigach, Jakiw. 2022. Score-based generative models detect manifolds. Pages 35852–35865 of: Koyejo, S., Mohamed, S., Agarwal, A., Belgrave, D., Cho, K., and Oh, A. (eds), *Advances in Neural Information Processing Systems*, vol. 35. Curran Associates, Inc.
- Polyak, Boris T. 1963. Gradient methods for minimizing functionals. *Ž. Vyčisl. Mat i Mat. Fiz.*, **3**, 643–653.
- Pooladian, Aram-Alexandre, Ben-Hamu, Heli, Domingo-Enrich, Carles, Amos, Brandon, Lipman, Yaron, and Chen, Ricky T. Q. 2023. Multisample flow matching: straightening flows with minibatch couplings. Pages 28100–28127 of: Krause, Andreas, Brunskill, Emma, Cho, Kyunghyun, Engelhardt, Barbara, Sabato, Sivan, and Scarlett, Jonathan (eds), *Proceedings of the 40th International Conference on Machine Learning*. Proceedings of Machine Learning Research, vol. 202. PMLR.
- Potapchik, Peter, Azangulov, Iskander, and Deligiannidis, George. 2025. Linear convergence of diffusion models under the manifold hypothesis. Pages 4668–4685 of: Haghtalab, Nika, and Moitra, Ankur (eds), *Proceedings of Thirty Eighth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 291. PMLR.
- Rademacher, Luis, and Vempala, Santosh. 2008. Dispersion of mass and the complexity of randomized geometric algorithms. *Adv. Math.*, **219**(3), 1037–1069.
- Raginsky, Maxim, and Rakhlin, Alexander. 2011. Information-based complexity, feedback and dynamics in convex programming. *IEEE Trans. Inform. Theory*, **57**(10), 7036–7056.

- Raginsky, Maxim, Rakhlin, Alexander, and Telgarsky, Matus. 2017. Non-convex learning via stochastic gradient Langevin dynamics: a non-asymptotic analysis. Pages 1674–1703 of: Kale, Satyen, and Shamir, Ohad (eds), *Proceedings of the 2017 Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 65. PMLR.
- Rassoul-Agha, Firas, and Seppäläinen, Timo. 2015. *A course on large deviations with an introduction to Gibbs measures*. Graduate Studies in Mathematics, vol. 162. American Mathematical Society, Providence, RI.
- Reeves, Galen, and Pfister, Henry D. 2025. Information-theoretic proofs for diffusion sampling. Pages 1–6 of: 2025 *IEEE International Symposium on Information Theory (ISIT)*.
- Roberts, Gareth O., and Rosenthal, Jeffrey S. 1998. Optimal scaling of discrete approximations to Langevin diffusions. *J. R. Stat. Soc. Ser. B Stat. Methodol.*, **60**(1), 255–268.
- Roberts, Gareth O., and Tweedie, Richard L. 1996. Exponential convergence of Langevin distributions and their discrete approximations. *Bernoulli*, **2**(4), 341–363.
- Rockafellar, R. Tyrrell. 1966. Characterization of the subdifferentials of convex functions. *Pacific J. Math.*, **17**, 497–510.
- Rothaus, Oscar S. 1985. Analytic inequalities, isoperimetric inequalities and logarithmic Sobolev inequalities. *J. Funct. Anal.*, **64**(2), 296–313.
- Roussel, Julien, and Stoltz, Gabriel. 2018. Spectral methods for Langevin dynamics and associated error estimates. *ESAIM. Mathematical Modelling and Numerical Analysis*, **52**(3), 1051–1083.
- Sahin, M. Berk, Sharif, Behzad, and Hashemi, Abolfazl. 2026. Zeroth-order non-log-concave sampling with variance reduction and applications to inverse problems. *arXiv preprint 2605.30573*.
- Salez, Justin, and Youssef, Pierre. 2025. Intrinsic regularity in the discrete log-Sobolev inequality. *arXiv preprint 2503.02793*.
- Salez, Justin, Tikhomirov, Konstantin, and Youssef, Pierre. 2023. Upgrading MLSI to LSI for reversible Markov chains. *J. Funct. Anal.*, **285**(9), Paper No. 110076, 15.
- Salim, Adil, Kovalev, Dmitry, and Richtarik, Peter. 2019. Stochastic proximal Langevin algorithm: potential splitting and nonasymptotic rates. In: Wallach, H., Larochelle, H., Beygelzimer, A., d'Alché-Buc, F., Fox, E., and Garnett, R. (eds), *Advances in Neural Information Processing Systems*, vol. 32. Curran Associates, Inc.
- Salim, Adil, Korba, Anna, and Luise, Giulia. 2020. The Wasserstein proximal gradient algorithm. Pages 12356–12366 of: Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M.F., and Lin, H. (eds), *Advances in Neural Information Processing Systems*, vol. 33. Curran Associates, Inc.
- Santambrogio, Filippo. 2015. *Optimal transport for applied mathematicians*. Progress in Nonlinear Differential Equations and their Applications, vol. 87. Birkhäuser/Springer, Cham. Calculus of variations, PDEs, and modeling.
- Sanz-Serna, Jesús María, and Zygalakis, Konstantinos C. 2021. Wasserstein distance estimates for the distributions of numerical approximations to ergodic stochastic differential equations. *Journal of Machine Learning Research*, **22**(242), 1–37.
- Saumard, Adrien, and Wellner, Jon A. 2014. Log-concavity and strong log-concavity: a review. *Stat. Surv.*, **8**, 45–114.
- Schlichting, André. 2019. Poincaré and log–Sobolev inequalities for mixtures. *Entropy*, **21**(1).
- Schuh, Katharina. 2024. Global contractivity for Langevin dynamics with distribution-dependent forces and uniform in time propagation of chaos. *Annales de l'Institut Henri Poincaré, Probabilités et Statistiques*, **60**(2), 753–789.
- Shen, Lingqing, Balasubramanian, Krishnakumar, and Ghadimi, Saeed. 2019. Non-asymptotic results for Langevin Monte Carlo: coordinate-wise and black-box sampling. *arXiv preprint 1902.01373v3*.
- Shen, Ruoqi, and Lee, Yin Tat. 2019. The randomized midpoint method for log-concave sampling. In: Wallach, H., Larochelle, H., Beygelzimer, A., d'Alché-Buc, F., Fox, E., and Garnett, R. (eds), *Advances in Neural Information Processing Systems*, vol. 32. Curran Associates, Inc.
- Shi, Bobby, Tian, Kevin, and Zhang, Matthew S. 2025. Perspectives on stochastic localization. *arXiv preprint 2510.04460*.
- Shih, Andy, Belkhale, Suneel, Ermon, Stefano, Sadigh, Dorsa, and Anari, Nima. 2023. Parallel sampling of diffusion models. Pages 4263–4276 of: Oh, A., Naumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds), *Advances in Neural Information Processing Systems*, vol. 36. Curran Associates, Inc.
- Sohl-Dickstein, Jascha, Weiss, Eric, Maheswaranathan, Niru, and Ganguli, Surya. 2015. Deep unsupervised learning using nonequilibrium thermodynamics. Pages 2256–2265 of: Bach, Francis, and Blei, David (eds), *Proceedings of the 32nd International Conference on Machine Learning*. Proceedings of Machine Learning Research, vol. 37. Lille, France: PMLR.

- Song, Jiaming, Meng, Chenlin, and Ermon, Stefano. 2021a. Denoising diffusion implicit models. In: *International Conference on Learning Representations*.
- Song, Yang, and Ermon, Stefano. 2019. Generative modeling by estimating gradients of the data distribution. In: Wallach, H., Larochelle, H., Beygelzimer, A., d'Alché-Buc, F., Fox, E., and Garnett, R. (eds), *Advances in Neural Information Processing Systems*, vol. 32. Curran Associates, Inc.
- Song, Yang, Sohl-Dickstein, Jascha, Kingma, Diederik P., Kumar, Abhishek, Ermon, Stefano, and Poole, Ben. 2021b. Score-based generative modeling through stochastic differential equations. In: *International Conference on Learning Representations*.
- Srinivasan, Rishikesh, and Nagaraj, Dheeraj. 2025. Poisson midpoint method for log concave sampling: beyond the strong error lower bounds. *arXiv preprint 2506.07614*.
- Steele, J. Michael. 2001. *Stochastic calculus and financial applications*. Applications of Mathematics (New York), vol. 45. Springer-Verlag, New York.
- Sturm, Karl-Theodor. 2006a. On the geometry of metric measure spaces. I. *Acta Math.*, **196**(1), 65–131.
- Sturm, Karl-Theodor. 2006b. On the geometry of metric measure spaces. II. *Acta Math.*, **196**(1), 133–177.
- Su, Weijie, Boyd, Stephen, and Candès, Emmanuel J. 2016. A differential equation for modeling Nesterov's accelerated gradient method: theory and insights. *J. Mach. Learn. Res.*, **17**, Paper No. 153, 43.
- Sudakov, Vladimir N., and Cirel'son, Boris S. 1974. Extremal properties of half-spaces for spherically invariant measures. *Zap. Naučn. Sem. Leningrad. Otdel. Mat. Inst. Steklov. (LOMI)*, **41**, 14–24, 165. Problems in the theory of probability distributions, II.
- Talagrand, Michel. 1991. A new isoperimetric inequality and the concentration of measure phenomenon. Pages 94–124 of: *Geometric aspects of functional analysis (1989–90)*. Lecture Notes in Math., vol. 1469. Springer, Berlin.
- Tong, Alexander, Fatras, Kilian, Malkin, Nikolay, Huguet, Guillaume, Zhang, Yanlei, Rector-Brooks, Jarrod, Wolf, Guy, and Bengio, Yoshua. 2024. Improving and generalizing flow-based generative models with minibatch optimal transport. *Transactions on Machine Learning Research*.
- van Erven, Tim, and Harremoës, Peter. 2014. Rényi divergence and Kullback–Leibler divergence. *IEEE Trans. Inform. Theory*, **60**(7), 3797–3820.
- van Handel, Ramon. 2016. *Probability in high dimension*.
- Vargas, Francisco, Thodoroff, Pierre, Lamacraft, Austen, and Lawrence, Neil. 2021. Solving Schrödinger bridges via maximum likelihood. *Entropy*, **23**(9), Paper No. 1134, 30.
- Vempala, Santosh, and Wibisono, Andre. 2019. Rapid convergence of the unadjusted Langevin algorithm: isoperimetry suffices. Pages 8094–8106 of: Wallach, H., Larochelle, H., Beygelzimer, A., d'Alché-Buc, F., Fox, E., and Garnett, R. (eds), *Advances in Neural Information Processing Systems 32*. Curran Associates, Inc.
- Vempala, Santosh S., and Wibisono, Andre. 2026. The geometry of efficient nonconvex sampling. *arXiv preprint 2603.25622*.
- Vershynin, Roman. 2018. *High-dimensional probability*. Cambridge Series in Statistical and Probabilistic Mathematics, vol. 47. Cambridge University Press, Cambridge. An introduction with applications in data science, With a foreword by Sara van de Geer.
- Villani, Cédric. 2003. *Topics in optimal transportation*. Graduate Studies in Mathematics, vol. 58. American Mathematical Society, Providence, RI.
- Villani, Cédric. 2009a. Hypocoercivity. *Mem. Amer. Math. Soc.*, **202**(950), iv+141.
- Villani, Cédric. 2009b. *Optimal transport*. Grundlehren der Mathematischen Wissenschaften [Fundamental Principles of Mathematical Sciences], vol. 338. Springer-Verlag, Berlin. Old and new.
- Vincent, Pascal. 2011. A connection between score matching and denoising autoencoders. *Neural Comput.*, **23**(7), 1661–1674.
- von Neumann, John. 1951. Various techniques used in connection with random digits. Chap. 13, pages 36–38 of: Householder, A. S., Forsythe, G. E., and Germond, H. H. (eds), *Monte Carlo method*. National Bureau of Standards Applied Mathematics Series, vol. 12. Washington, DC: US Government Printing Office.
- Wang, Feng-Yu. 1997. Logarithmic Sobolev inequalities on noncompact Riemannian manifolds. *Probab. Theory Related Fields*, **109**(3), 417–424.
- Wang, Feng-Yu. 2000. Functional inequalities for empty essential spectrum. *J. Funct. Anal.*, **170**(1), 219–245.
- Wang, Feng-Yu. 2011. Equivalent semigroup properties for the curvature-dimension condition. *Bull. Sci. Math.*, **135**(6–7), 803–815.

- Wang, Feng-Yu. 2013. *Harnack inequalities for stochastic partial differential equations*. SpringerBriefs in Mathematics. Springer, New York.
- Wang, Feng-Yu. 2014. *Analysis for diffusion processes on Riemannian manifolds*. Advanced Series on Statistical Science & Applied Probability, vol. 18. World Scientific Publishing Co. Pte. Ltd., Hackensack, NJ.
- Welling, Max, and Teh, Yee-Whye. 2011. Bayesian learning via stochastic gradient Langevin dynamics. Pages 681–688 of: *Proceedings of the 28th International Conference on Machine Learning*. ICML'11. Madison, WI, USA: Omnipress.
- Wibisono, Andre. 2018. Sampling as optimization in the space of measures: the Langevin dynamics as a composite optimization problem. Pages 2093–3027 of: Bubeck, Sébastien, Perchet, Vianney, and Rigollet, Philippe (eds), *Proceedings of the 31st Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 75. PMLR.
- Wibisono, Andre. 2025. Mixing time of the proximal sampler in relative Fisher information via strong data processing inequality (extended abstract). Pages 5716–5717 of: Haghtalab, Nika, and Moitra, Ankur (eds), *Proceedings of Thirty Eighth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 291. PMLR.
- Wu, Keru, Schmidler, Scott, and Chen, Yuansi. 2021. Minimax mixing time of the Metropolis-adjusted Langevin algorithm for log-concave sampling. *arXiv e-prints*.
- Xun, Zhiyang, and Price, Eric. 2026. Query lower bounds for diffusion sampling. *arXiv preprint 2604.10857*.
- Yu, Lu, Karagulyan, Avetik, and Dalalyan, Arnak S. 2024. Langevin Monte Carlo for strongly log-concave distributions: randomized midpoint revisited. In: *The Twelfth International Conference on Learning Representations*.
- Zhang, Kelvin S., Peyré, Gabriel, Fadili, Jalal, and Pereyra, Marcelo. 2020. Wasserstein control of mirror Langevin Monte Carlo. Pages 3814–3841 of: Abernethy, Jacob, and Agarwal, Shivani (eds), *Proceedings of Thirty Third Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 125. PMLR.
- Zhang, Matthew S. 2025. Analysis of Langevin midpoint methods using an anticipative Girsanov theorem. *arXiv preprint 2507.12791*.
- Zhang, Matthew S., Chewi, Sinho, Li, Mufan (Bill), Balasubramanian, Krishnakumar, and Erdogdu, Murat A. 2023. Improved discretization analysis for underdamped Langevin Monte Carlo. Pages 36–71 of: Neu, Gergely, and Rosasco, Lorenzo (eds), *Proceedings of Thirty Sixth Conference on Learning Theory*. Proceedings of Machine Learning Research, vol. 195. PMLR.
- Zhang, Matthew S., Huan, Stephen, Huang, Jerry, Boffi, Nicholas M., Chen, Sitan, and Chewi, Sinho. 2025. Sublinear iterations can suffice even for DDPMs. *arXiv preprint 2511.04844*.
- Zhang, Matthew S., Altschuler, Jason M., and Chewi, Sinho. 2026. Algorithmic warm starts for Hamiltonian Monte Carlo. *arXiv preprint 2603.22741*.
- Zhou, Huanjian, and Sugiyama, Masashi. 2025a. The adaptive complexity of minimizing relative Fisher information. Pages 76789–76812 of: Belgrave, D., Zhang, C., Lin, H., Pascanu, R., Koniusz, P., Ghassemi, M., and Chen, N. (eds), *Advances in Neural Information Processing Systems*, vol. 38. Curran Associates, Inc.
- Zhou, Huanjian, and Sugiyama, Masashi. 2025b. Parallel simulation for log-concave sampling and score-based diffusion models. In: *Forty-second International Conference on Machine Learning*.
- Zimmermann, David. 2013. Logarithmic Sobolev inequalities for mollified compactly supported measures. *J. Funct. Anal.*, **265**(6), 1064–1083.
- Zimmermann, David. 2016. Elementary proof of logarithmic Sobolev inequalities for Gaussian convolutions on $\mathbb{R}$. *Ann. Math. Blaise Pascal*, **23**(1), 129–140.
- Zou, Difan, and Gu, Quanquan. 2021. On the convergence of Hamiltonian Monte Carlo with stochastic gradients. Pages 13012–13022 of: Meila, Marina, and Zhang, Tong (eds), *Proceedings of the 38th International Conference on Machine Learning*. Proceedings of Machine Learning Research, vol. 139. PMLR.
- Zou, Difan, Xu, Pan, and Gu, Quanquan. 2018. Subsampled stochastic variance-reduced gradient Langevin dynamics. Pages 508–518 of: *34th Conference on Uncertainty in Artificial Intelligence*, vol. 1.

- c -concavity, 44
- c -conjugate potentials, 43
- f -divergence, 39, 218
- absolute continuity, 26
- abstract Wiener space, 102
- acceleration, 150
- Bakry–Émery theorem, 19
- Benamou–Brenier formula, 44
- blow-up, 68
- Bochner identity, 50
- Bonnet–Myers diameter bound, 83
- Brascamp–Lieb inequality, 52, 259
- Bregman divergence, 53, 260
- Bregman proximal lemma, 262
- Bregman transport cost, 53, 263
- Bregman transport inequality, 53, 88, 264
- Brenier potential, 24
- Brenier’s theorem, 22
 - Bregman transport costs, 263
- Brownian bridge, 106
- Brownian motion, 4
- Brunn–Minkowski inequality, 91
- Caffarelli’s contraction theorem, 61, 113
- Cameron–Martin
 - space, 101
 - theorem, 101
- carré du champ, 14
 - iterated operator, 18, 50
- Cheeger’s inequality, 199
- chi-squared divergence, 17
- coarea inequality, 73
- cold start, 195
- concentration function, 68
- conductance, 199
 - s -conductance, 202
- continuity equation, 26
- covariant derivative, 78
- Cramér–Rao inequality, 113
- curvature
 - Gaussian, 80
 - Ricci, 80
 - coarse, 85
 - Riemann, 80
 - scalar, 80
 - sectional, 80
- curvature-dimension condition, 19, 50, 82
- data-processing inequality, 39
- descent lemma, 128
- diffusion coefficient, 8
- diffusion semigroup, 19, 54
- Dirichlet energy, 14, 198
- displacement interpolation, 30
- divergence, 79
- divergence lemma, 163
- Donsker–Varadhan variational principle, 40
- Doob martingale, 41
- Doob’s transform, 104
- drift coefficient, 8
- Efron–Stein inequality, 41
- elementary process, 5
- entropy, 238
- Euler–Maruyama discretization, 123
- evolution variational inequality (EVI), 132
- exponential map, 30
- Eyring–Kramers formula, 59
- Föllmer drift, 108
- Fano’s inequality, 238
- finite variation, 96
- first variation, 31
- Fisher information, 18, 55
- flow map, 146
- Fokker–Planck equation, 12
- friction, 149
- generalized geodesic, 264
- geodesic, 29, 79
- geodesic convexity, 31

- Gibbs sampling, 255
- Girsanov's theorem, 104, 134
 - shifted Girsanov, 116, 135
- Grönwall's lemma, 16
- gradient, 78
- gradient descent, 124
 - accelerated, 150
- Hörmander's L^2 method, 52
- Hamilton's equations of motion, 146
- Hamiltonian, 146
- Hamiltonian Monte Carlo (HMC)
 - ideal, 145
 - Metropolized (MHMC), 193
- Hamilton–Jacobi equation, 44
- Harnack inequality, 36, 47, 92, 116, 139
 - log-Harnack, 92
- Helffer–Sjöstrand identity, 89
- Hellinger distance, 278
- Holley–Stroock perturbation, 60
- Hopf–Lax semigroup, 45
- hypercontractivity, 90, 177
- hypocoercivity, 155
 - entropic, 158
- infinitesimal generator, 11, 198
- isoperimetric profile, 71
- isoperimetry
 - Cheeger, 72, 201
 - Gaussian, 71, 77, 205
 - sphere, 70, 93
- Itô integral, 3
- Itô isometry, 6
- Itô process, 8
- Itô's formula, 9, 99
- Kannan–Lovász–Simonovits (KLS) conjecture, 59
- Kantorovich problem, 20
- Kazamaki's condition, 104
- Kim–Milman map, 117
- Knothe–Rosenblatt map, 256
- Kolmogorov's backward equation, 12
- Kolmogorov's forward equation, 12
- Kullback–Leibler (KL) divergence, 17
 - chain rule, 40
- Langevin diffusion, 3
- Langevin Monte Carlo (LMC), 123
- Laplace–Beltrami operator, 79
- Latala–Oleszkiewicz inequality (LOI), 86
- lazy chain, 195
- Le Cam's inequality, 278
- leapfrog integrator, 193
- Levi–Civita connection, 78
- Lichnerowicz inequality, 94
- Lie bracket, 80
- local error analysis, 142
- local martingale, 7
- localization, 6
- localizing sequence, 6
- log-Sobolev inequality (LSI), 18, 48, 128
 - defective, 64
 - local, 56
 - modified, 84, 198
- logarithmic map, 30
- Malliavin calculus, 115
- manifold, 29
 - Hessian, 78
- Markov semigroup, 10
- martingale, 5
 - exponential, 103
- Marton's tensorization, 67
- McCann's interpolation, 30
- mean-squared error analysis, *see* local error analysis
- mesh, 96
- metastability, 273
- metric derivative, 26
- Metropolis-adjusted Langevin algorithm (MALA), 193
- Metropolis–Hastings (MH) filter, 191
- Metropolized random walk (MRW), 193
- Minkowski content, 71
- mirror descent, 257
- mirror Langevin, 257
- mirror Langevin Monte Carlo (MLMC), 261
- mirror map, 257
- Monge problem, 20
- Monge–Ampère equation, 87
- mutual information, 238
- Newton–Langevin diffusion, 259
- no-U-turn sampler (NUTS), 166
- Novikov's condition, 104
- optimal transport, 20
 - dual, 21
 - entropic regularization, 111
 - fundamental theorem, 22
 - optimal transport plan, 20
- Ornstein–Uhlenbeck (OU) process, 42, 155
- Otto calculus, 33
- Otto–Villani theorem, 36
- parallel transport, 78
- Pinsker's inequality, 37
- Poincaré inequality (PI), 16, 48, 198
 - L^p – L^q , 74
 - local, 55
 - mirror, 260
- Poisson bracket, 169
- Poisson equation, 52, 87
- Polyak–Łojasiewicz (PL) inequality, 35, 59, 128, 231
- Prékopa–Leindler inequality, 90
- proximal oracle, 214
- proximal point method (PPM), 214
- proximal sampler, 215
- quadratic variation, 97
- Rényi divergence, 57, 174
- randomized midpoint discretization, 141
- reflection coupling, 138
- rejection sampling, 189
- relative convexity, 260
- relative smoothness, 260

- restricted Gaussian oracle (RGO), 215
- reverse transport inequality, *see* Harnack inequality
- reversibility, 13, 191
- reweighted functional inequality, 278
- Ricci curvature, 51
- Riemannian metric, 29, 77
- Schrödinger bridge, 110
- score matching
 - implicit, 284
- second-order lift, 161
- self-concordance, 262
- semimartingale, 98
- shifted ODE algorithm, 166
- Sobolev inequality, 83
- Sobolev norm, 88
 - inverse, 88
- stochastic calculus, 3
- stochastic differential equation (SDE), 9
- stochastic gradient Langevin Monte Carlo (SG-LMC), 249
- stochastic localization, 109
- stopping time, 6
- symplectic integrator, 193
- synchronous coupling, 138
- Talagrand's T_1 inequality, 48
- Talagrand's T_2 inequality, 36, 48
- tangent bundle, 77
- tangent space, 29, 77
- time reversal, 106, 117, 220
- total variation, 96
 - total variation (TV) distance, 38, 97
 - total variation norm, 96
- Tweedie's identity, 288
- underdamped Langevin diffusion, 149
- vector field, 78
- volume measure, 79
- warm start, 195
- Wasserstein gradient flow, 32
- Wasserstein metric, 20, 44
- weak triangle inequality, 181
- Wiener measure, 100