

Split-Flows: Measure Transport and Information Loss Across Molecular Resolutions

Sander Hummerich
Institute for Theoretical Physics
Heidelberg University

Tristan Bereau
Institute for Theoretical Physics
Heidelberg University

Ullrich Köthe
Computer Vision and Learning Lab
Heidelberg University

Abstract

By reducing resolution, coarse-grained models greatly accelerate molecular simulations, unlocking access to long-timescale phenomena, though at the expense of microscopic information. Recovering this fine-grained detail is essential for tasks that depend on atomistic accuracy, making backmapping a central challenge in molecular modeling. We introduce split-flows, a novel flow-based approach that reinterprets backmapping as a continuous-time measure transport across resolutions. Unlike existing generative strategies, split-flows establish a direct probabilistic link between resolutions, enabling expressive conditional sampling of atomistic structures and—for the first time—a tractable route to computing mapping entropies, an information-theoretic measure of the irreducible detail lost in coarse-graining. We demonstrate these capabilities on diverse molecular systems, including chignolin, a lipid bilayer, and alanine dipeptide, highlighting split-flows as a principled framework for accurate backmapping and systematic evaluation of coarse-grained models. Our code is available at <https://github.com/BereauLab/split-flows>.

1 INTRODUCTION

Coarse-grained models play a central role in molecular and material simulations [Noid, 2023]. By marginalizing out unnecessary detail, they drastically reduce the computational cost of simulation and smooth out

Proceedings of the 29th International Conference on Artificial Intelligence and Statistics (AISTATS) 2026, Tangier, Morocco. PMLR: Volume 300. Copyright 2026 by the author(s).

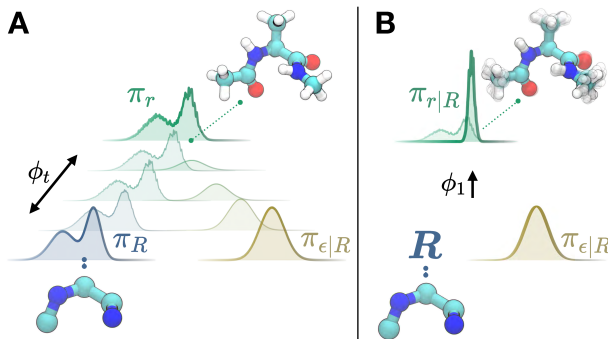


Figure 1: (A) Split-flows connect fine- and coarse-grained densities, π_r and π_R , respectively, at different molecular resolutions via a continuous-time measure transport that maps the excess degrees of freedom of the fine-grained resolution to a simple noise distribution, $\pi_{\epsilon|R}$. (B) This enables sampling from the conditional density $\pi_{r|R}$, i.e., generative backmapping, and quantifies the information loss inherent in the coarse-grained representation.

the underlying energy landscape. This enables simulations on length and time scales that are otherwise intractable in fine-grained models, providing an efficient tool to study slow collective dynamics and mesoscale phenomena such as protein folding, polymer conformational transitions, membrane remodeling, lipid-domain organization, and large-scale self-assembly.

A coarse-graining map implicitly defines an ill-posed inverse problem referred to as backmapping; that is, to reconstruct the marginalized degrees of freedom of the fine-grained model from the coarse-grained representation. As the forward process defines a many-to-one map—many detailed configurations are mapped to the same coarse-grained configuration—the reverse process can be cast as a generative-modeling problem: learning a probabilistic model for the distribution of fine-grained configurations corresponding to each coarse-grained representative.

The reduction of degrees of freedom in coarse-grained

models inevitably leads to information loss relative to the fine-grained description. This loss can be quantified through the concept of mapping entropy [Shell, 2008, Foley et al., 2015], which measures the average entropy of the distribution of fine-grained configurations that map to a given coarse-grained representative. Mapping entropy thus provides an information-theoretic lens on multiscale modeling, where a low mapping entropy corresponds to high information loss due to reduced resolution. This perspective enables quantitative assessment of the information loss in coarse-grained models and can ultimately inform model and simulation design.

In this work, we propose *split-flows*—a novel flow-based model that provides a clear approach to bridging the dimensional gap between fine- and coarse-grained domains, as illustrated in Figure 1. Split-flows define a continuous-time measure transport across dimensions, enabling us to connect the configurational densities at two different resolutions for general coarse-graining strategies. In addition to addressing the backmapping problem, this probabilistic link between fine- and coarse-grained resolutions allows us to compute the information loss of the coarse-graining map. In summary, we make the following contributions:

- **Method:** We introduce split-flows, a flow-based model that enables continuous-time transport of probability measures across different resolutions, bridging fine- and coarse-grained domains.
- **Theory:** We show that split-flows allow, for the first time, tractable and general computation of mapping entropy for arbitrary coarse-graining maps, providing a principled measure of information loss.
- **Applications:** We apply split-flows to diverse biomolecular systems—chignolin, a lipid bilayer, and alanine dipeptide—demonstrating accurate backmapping and their utility for information-theoretic assessment of coarse-grained models.

2 RELATED WORK

Solving the inverse problem of backmapping is a central challenge in multiscale molecular modeling [Peter and Kremer, 2009]. Mirroring trends across many scientific domains, data-driven methods increasingly replace traditional handcrafted algorithms, such as those by [Rzepiela et al., 2010], and [Wassenaar et al., 2014], which predict approximate fine-grained configurations from coarse inputs, followed by costly refinement. Early approaches, such as [Stieffenhofer et al., 2020], [Li et al., 2020], and

[Wang et al., 2022], leverage generative adversarial networks and variational autoencoders to generate fine-grained samples, without the need for post hoc refinement. [Shmilovich et al., 2022] extend this line of work by incorporating information along reconstructed simulations to ensure temporal consistency. More recent methods by [Jones et al., 2023], [Jones et al., 2025, Berлага et al., 2025], and [Torre and Sugita, 2025] adopt multi-step samplers, i.e., continuous normalizing flows and diffusion models, enabling generalization to unseen structures through residue-wise processing and transferable coarse-graining schemes. While these models emphasize energetic plausibility, transferability, or dynamical consistency, they do not establish a probabilistic link between resolutions and therefore miss key statistical properties of the coarse-graining map. Our method addresses this limitation.

Normalizing flows, introduced by [Rezende and Mohamed, 2015] in discrete form, map complex data distributions to simple latents. The continuous-time formulation of [Chen et al., 2018] improves expressiveness but initially lacks a tractable training procedure. Flow matching [Lipman et al., 2023] resolves this by replacing maximum likelihood with a quadratic regression objective for the underlying velocity field, enabling efficient training of continuous normalizing flows. Modern formulations of flow matching, particularly those by [Albergo et al., 2023] and [Tong et al., 2024], generalize normalizing flows to define a measure transport between arbitrary pairs of distributions. Most similar to our approach, [Albergo et al., 2024] apply this framework to image super-resolution and in-painting. We build on this modern interpretation of continuous normalizing flows to connect molecular configurations across resolutions.

Mapping entropy, introduced by [Shell, 2008] and subsequently formalized by [Rudzinski and Noid, 2011, Foley et al., 2015], quantifies the information lost when a fine-grained system is represented at reduced resolution. As such, it provides a principled criterion for analyzing coarse-grained representations. Prior work has used mapping entropy to characterize the entropic structure of coarse-grained models [Kidder et al., 2021], to identify informative mappings for specific systems such as actin [Kidder and Noid, 2024], and to study the extent to which structural reduction preserves dynamical behavior across resolutions [Armstrong et al., 2012, Jin et al., 2023]. It has also been used to disentangle entropic and energetic contributions to collective variables [Mussi and Noid, 2025], implemented in practical coarse-graining software frame-

works [Giulini et al., 2020, Giulini et al., 2024], and applied beyond molecular modeling to settings such as spin systems and low-dimensional financial descriptors [Holtzman et al., 2022]. Despite this breadth of applications, existing approaches are often tailored to particular models or system classes. By contrast, split-flows provide a general and rigorous framework for computing mapping entropies across a broad class of systems and reduction strategies.

3 PRELIMINARIES

3.1 Thermodynamic Framework

Notation: We use lowercase variable names to denote quantities at the fine-grained resolution and uppercase variable names for quantities at the coarse-grained resolution. For ease of presentation, we treat variables as unitless quantities.

We consider a system with n degrees of freedom at temperature T and configurations denoted by $\mathbf{r}$. At equilibrium, these configurations follow the Boltzmann distribution governed by the potential energy function $u : \mathbb{R}^n \rightarrow \mathbb{R}$:

$$\pi_{\mathbf{r}}(\mathbf{r}) = Z^{-1} \exp[-u(\mathbf{r})/(k_B T)], \quad (1)$$

where $Z = \int_{\mathbb{R}^n} d\mathbf{r} \exp[-u(\mathbf{r})/(k_B T)]$ is the normalization constant and k_B is the Boltzmann constant.

In practice, samples from $\pi_{\mathbf{r}}$ are generated using trajectory-based methods such as molecular dynamics or Monte Carlo simulations. These methods often suffer from slow convergence at the fine-grained resolution, since the energy landscape is rugged and trajectories can become trapped in local minima. Coarse-grained models accelerate sampling by both reducing the number of degrees of freedom and smoothing out the underlying energy surface [Noid, 2013].

3.2 Coarse-Graining

In this work, we focus on so-called bottom-up coarse-graining approaches that derive a coarse representation from a fine-grained model via a coarse-graining map

$$M : \mathbb{R}^n \rightarrow \mathbb{R}^N, \quad \mathbf{R} = M(\mathbf{r}), \quad (2)$$

which assigns to each fine-grained configuration $\mathbf{r}$ a coarse-grained representative $\mathbf{R}$ with N degrees of freedom, as shown in Figure 2. Such maps aim to preserve the essential physics, effectively separating *slow* (typically complex) from *fast* (typically simple) degrees of freedom.

The corresponding coarse-grained density $\pi_{\mathbf{R}}$ is ob-

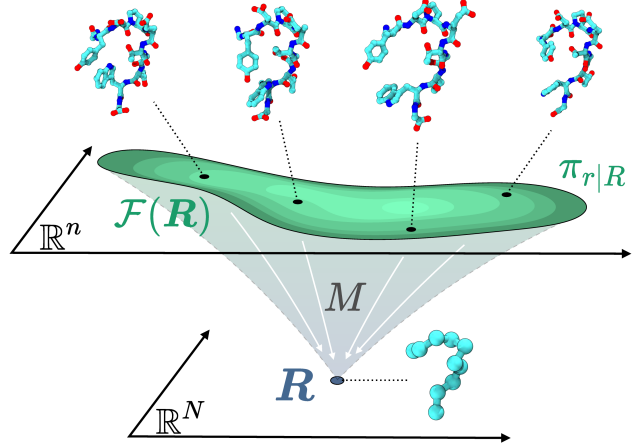


Figure 2: Bottom-up coarse-graining defines a many-to-one mapping operator M that reduces a set $\mathcal{F}(\mathbf{R})$ of fine-grained configurations to a single coarse-grained representative $\mathbf{R}$.

tained by integrating out the fast degrees of freedom:

$$\pi_{\mathbf{R}}(\mathbf{R}) = \int_{\mathbb{R}^n} d\mathbf{r} \delta(M(\mathbf{r}) - \mathbf{R}) \pi_{\mathbf{r}}(\mathbf{r}), \quad (3)$$

where the delta function ensures that only fine-grained configurations consistent with $\mathbf{R}$ contribute. Since this exact density is generally intractable, coarse-graining methods approximate $\pi_{\mathbf{R}}$ with a model $\hat{\pi}_{\mathbf{R}}$ that ideally preserves consistency with the above equation.

In this work, however, we assume access to the exact coarse-grained distribution. Analogous to the fine-grained Boltzmann distribution, it can be written as

$$\pi_{\mathbf{R}}(\mathbf{R}) \propto \exp[-W(\mathbf{R})/(k_B T)]. \quad (4)$$

Here, the effective potential $W : \mathbb{R}^N \rightarrow \mathbb{R}$ is a *free energy*,

$$W(\mathbf{R}) = E(\mathbf{R}) - TS(\mathbf{R}), \quad (5)$$

which includes energetic and entropic contributions [Noid, 2013]; $E(\mathbf{R})$ is the mean fine-grained energy conditioned on $\mathbf{R}$, and $S(\mathbf{R})$ is an entropic term, that measures the structure of the distribution of compatible fine-grained configurations. Coarse-graining averages over microscopic energies while introducing an entropic bias toward states with many realizations. This results in a smoother free-energy landscape $W(\mathbf{R})$ that is easier to sample than the atomistic potential, at the cost of information loss, as different fine-grained configurations mapping to the same $\mathbf{R}$ become indistinguishable.

3.3 Information Loss in Coarse-Grained Representations

A quantitative measure of information loss in coarse-grained representations can be derived from the con-

cept of mapping entropy S_{map} and its configuration-dependent (*local*) counterpart $S(\mathbf{R})$.

To introduce the mapping entropy, we first define the *fiber* associated with a coarse-grained representative. The fiber is the pre-image of $\mathbf{R}$ under the mapping M , i.e., it is the *lost subensemble* [Kidder and Noid, 2024] of all fine-grained states that map to $\mathbf{R}$:

$$\mathcal{F}(\mathbf{R}) = \{\mathbf{r} \in \mathbb{R}^n \mid M(\mathbf{r}) = \mathbf{R}\}. \quad (6)$$

Bayes’ theorem gives the *fiber distribution*—the conditional probability of a fine-grained configuration given its coarse-grained representative:

$$\pi_{\mathbf{r}|\mathbf{R}}(\mathbf{r} \mid \mathbf{R}) = \frac{\pi_{\mathbf{r}}(\mathbf{r})}{\pi_{\mathbf{R}}(\mathbf{R})}, \forall \mathbf{r} \in \mathcal{F}(\mathbf{R}). \quad (7)$$

We will denote the expectation of some d -dimensional observable $O : \mathbb{R}^n \rightarrow \mathbb{R}^d$ on the fine-grained configuration space as the *fiber average*:

$$\mathbb{E}_{\mathbf{r}|\mathbf{R}}[O(\mathbf{r})] = \int_{\mathcal{F}(\mathbf{R})} d\mathbf{r} \pi_{\mathbf{r}|\mathbf{R}}(\mathbf{r} \mid \mathbf{R}) O(\mathbf{r}) \quad (8)$$

which lets us evaluate observables over fine-grained states consistent with one particular coarse-grained representative, e.g., the energetic component $E(\mathbf{R}) = \mathbb{E}_{\mathbf{r}|\mathbf{R}}[u(\mathbf{r})]$ in Equation 5.

Using Equation 7, we can write the entropy of the fiber distribution as

$$\begin{aligned} S(\mathbf{R}) &= -k_B \int_{\mathcal{F}(\mathbf{R})} d\mathbf{r} \pi_{\mathbf{r}|\mathbf{R}}(\mathbf{r} \mid \mathbf{R}) \log \pi_{\mathbf{r}|\mathbf{R}}(\mathbf{r} \mid \mathbf{R}) \\ &= -k_B \mathbb{E}_{\mathbf{r}|\mathbf{R}} \left[\log \frac{\pi_{\mathbf{r}}(\mathbf{r})}{\pi_{\mathbf{R}}(\mathbf{R})} \right], \end{aligned} \quad (9)$$

which we denote the *local mapping entropy*. As outlined in Appendix A.1, this is the entropic contribution $S(\mathbf{R})$ in Equation 5. For compact domains, e.g., a periodic box, we can define the local excess mapping entropy as the relative entropy of the fiber distribution compared to the best guess we can make without any prior information, i.e., a uniform distribution over $\mathcal{F}(\mathbf{R})$:

$$S_e(\mathbf{R}) = -k_B \mathbb{E}_{\mathbf{r}|\mathbf{R}} \left[\log \frac{\pi_{\mathbf{r}|\mathbf{R}}(\mathbf{r} \mid \mathbf{R})}{u_{\mathbf{r}|\mathbf{R}}(\mathbf{r} \mid \mathbf{R})} \right] \quad (10)$$

which is the Kullback-Leibler divergence between the fiber distribution $\pi_{\mathbf{r}|\mathbf{R}}$ and a uniform distribution with density $u_{\mathbf{r}|\mathbf{R}}(\mathbf{r} \mid \mathbf{R}) = \text{Vol}(\mathcal{F}(\mathbf{R}))^{-1} \forall \mathbf{r} \in \mathcal{F}(\mathbf{R})$ defined over the fiber.

The local (excess) information loss due to reducing the fiber $\mathcal{F}(\mathbf{R})$ to a single representative $\mathbf{R}$ in the coarse-grained model relates to the local (excess) mapping entropy as

$$I_{(e)}(\mathbf{R}) = -S_e(\mathbf{R})/k_B. \quad (11)$$

It is evident that the local information loss $I(\mathbf{R})$ must be non-negative and thus $S(\mathbf{R}) \leq 0$. Taking the expectation of the local quantities $S(\mathbf{R})$ and $I(\mathbf{R})$ with respect to the coarse-grained density $\pi_{\mathbf{R}}$ then yields the global mapping entropy S_{map} and information loss I_{map} of the coarse-grained model.

3.4 Two-Sided Flow Matching

Two-sided flow matching aims to connect two non-trivial distributions π_0 and π_1 over an interpolation interval $[0, 1]$. Continuous normalizing flows (CNFs) define such a measure transport via the solution to an ordinary differential equation (ODE):

$$\frac{d}{dt} \phi_t(\mathbf{x}_0) = \mathbf{v}_t^\theta(\phi_t(\mathbf{x}_0)), \quad \phi_0(\mathbf{x}_0) = \mathbf{x}_0. \quad (12)$$

Here, $\mathbf{v}^\theta : [0, 1] \times \mathbb{R}^n \rightarrow \mathbb{R}^n$ is a time-dependent velocity field, which is parameterized by a neural network. The flow defines a continuous-time bijection between samples from the two endpoint distributions, π_0 and π_1 . The pushforward of the initial density π_0 under the flow ϕ_t is given by

$$\log \pi_t(\phi_t(\mathbf{x}_0)) = \log \pi_0(\mathbf{x}_0) - \int_0^t d\tau \nabla \cdot \mathbf{v}_\tau^\theta(\phi_\tau(\mathbf{x}_0)), \quad (13)$$

which defines a probability path between π_0 and π_1 .

Given a coupling $\pi_{0,1}$ of samples of two endpoint distributions π_0 and π_1 , [Albergo et al., 2023] propose the following quadratic regression objective:

$$\mathcal{L}_v(\theta) = \int_0^1 dt \mathbb{E}_{0,1} [\|\mathbf{v}_t^\theta(I_t(\mathbf{x}_0, \mathbf{x}_1)) - \partial_t I_t(\mathbf{x}_0, \mathbf{x}_1)\|^2], \quad (14)$$

which is a simple extension of the conditional flow matching objective, originally introduced by [Lipman et al., 2023]. The coupling $\pi_{0,1}$ defines how the flow should pair samples from the two endpoint distributions and is task-specific, e.g., an optimal transport coupling. It satisfies $\int_{\mathbb{R}^n} d\mathbf{x}_1 \pi_{0,1}(\mathbf{x}_0, \mathbf{x}_1) = \pi_0(\mathbf{x}_0)$ and $\int_{\mathbb{R}^n} d\mathbf{x}_0 \pi_{0,1}(\mathbf{x}_0, \mathbf{x}_1) = \pi_1(\mathbf{x}_1)$. The interpolant I_t is chosen to be of the form $I_t(\mathbf{x}_0, \mathbf{x}_1) = \alpha_t \mathbf{x}_0 + \beta_t \mathbf{x}_1$ and obeys the boundary conditions $\alpha_0 = \beta_1 = 1$ and $\alpha_1 = \beta_0 = 0$.

4 SPLIT-FLOWS

Notation: We identify the endpoint densities as $\pi_0 = \pi_{\mathbf{R}} \times \pi_{\epsilon|\mathbf{R}}$, an augmented coarse-grained density, and $\pi_1 = \pi_{\mathbf{r}}$, the fine-grained density. Correspondingly, $\mathbf{x}_0 = (\mathbf{R}, \epsilon)$ and $\mathbf{x}_1 = \mathbf{r}$. The augmented coarse-grained density $\pi_{\mathbf{R}} \times \pi_{\epsilon|\mathbf{R}}$ is introduced below.

Split-flows bridge the gap between distributions defined over domains with different dimensionality by

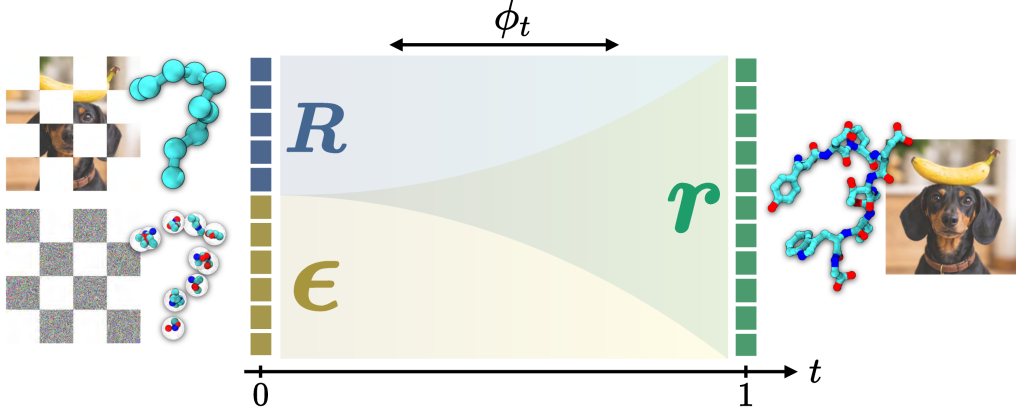


Figure 3: Split-flows define a one-to-one map between configurations of different resolutions. The lower-dimensional samples $\mathbf{R}$ are augmented with noise ϵ to resolve the degeneracy induced by the dimensionality gap. The flow ϕ_t connects the augmented coarse-grained configurations $(\mathbf{R}, \epsilon) \sim \pi_R \times \pi_{\epsilon|R}$ at $t = 0$ with the fine-grained configurations $\mathbf{r} \sim \pi_r$ at $t = 1$. An instructive analogy arises in image inpainting: a partially observed image is augmented with noise dimensions, and the split-flow acts as a probabilistic bridge to a complete, coherent one.

augmenting the lower-dimensional space with additional noise dimensions, as illustrated in Figure 3. Given the two endpoint distributions π_R and π_r defined over $\mathbb{R}^N$ and $\mathbb{R}^n$, respectively, we introduce a simple noise distribution $\pi_{\epsilon|R}$ on $\mathbb{R}^{n-N}$ and use a CNF ϕ_t , trained via the conditional flow matching objective in Equation 14, to learn a measure transport between $\pi_R \times \pi_{\epsilon|R}$ and π_r :

$$\phi_1 : \mathbb{R}^N \times \mathbb{R}^{n-N} \rightarrow \mathbb{R}^n, \quad (\mathbf{R}, \epsilon) \mapsto \phi_1(\mathbf{R}, \epsilon) = \mathbf{r}. \quad (15)$$

The noise distribution $\pi_{\epsilon|R}$ is chosen such that, given a coarse-grained representation, sampling is tractable, e.g., a Gaussian distribution. Using Equation 13 and the factorization of the augmented endpoint distribution, we can connect the densities π_R and π_r despite the difference in dimensionality via

$$\log \pi_r(\phi_1(\mathbf{R}, \epsilon)) = \log \pi_R(\mathbf{R}) + \log \pi_{\epsilon|R}(\epsilon | \mathbf{R}) - \int_0^1 d\tau \nabla \cdot \mathbf{v}_\tau^\theta(\phi_\tau(\mathbf{R}, \epsilon)). \quad (16)$$

In molecular modeling, we use this framework to relate the coarse-grained density π_R to the density over fine-grained configurations π_r . Introducing the conditional noise distribution $\pi_{\epsilon|R}$ resolves the many-to-one nature of coarse-graining and allows backmapping to be formulated as transport of measures across resolutions. From a geometric perspective, the flow learns a global coordinate transformation that disentangles the structure of fine-grained configuration space induced by the map M , i.e., its decomposition into slow and fast degrees of freedom. A more detailed discussion of this viewpoint is provided in Appendix A.5.

Algorithm 1 Per-sample loss computation

- 1: **Input:** fine-grained configuration $\mathbf{r}$, velocity field $\mathbf{v}^\theta$, coarse-graining map M , noise distribution $\pi_{\epsilon|R}$, interpolant I_t
- 2: Compute CG representation: $\mathbf{R} \leftarrow M(\mathbf{r})$
- 3: Sample noise: $\epsilon \sim \pi_{\epsilon|R}$
- 4: Sample time: $t \sim u_{[0,1]}$
- 5: Compute loss:

$$\mathcal{L}(\theta, \mathbf{r}) \leftarrow \|\mathbf{v}_t^\theta(I_t(\mathbf{R}, \epsilon, \mathbf{r})) - \partial_t I_t(\mathbf{R}, \epsilon, \mathbf{r})\|^2$$

- 6: **Output:** Per-sample loss $\mathcal{L}(\theta, \mathbf{r})$
-

To train split-flows in a two-sided manner, as outlined in Section 3.4, we pair samples from the two endpoint distributions using the coarse-graining map M , and construct a semi-deterministic coupling between $(\mathbf{R}, \epsilon)$ and $\mathbf{r}$:

$$\pi_{R, \epsilon, r}(\mathbf{R}, \epsilon, \mathbf{r}) = \pi_r(\mathbf{r}) \delta(\mathbf{R} - M(\mathbf{r})) \pi_{\epsilon|R}(\epsilon | \mathbf{R}). \quad (17)$$

This coupling encourages the flow to correctly pair fine-grained configurations with their respective coarse-grained counterparts and provides a straightforward way to evaluate a Monte Carlo estimate of the objective in Equation 14. We outline the per-sample loss computation in Algorithm 1.

This setup, once trained, allows us to easily access the fibers, i.e., the many possible fine-grained configurations mapping to a single coarse-grained representative, and the local mapping entropy of the coarse-graining map. We can generate samples $\mathbf{r} | \mathbf{R}$ from the conditional distribution $\pi_{r|R}$, i.e., samples on the

fiber $\mathcal{F}(\mathbf{R})$, using Algorithm 2.

Algorithm 2 Fiber-constrained sampling

- 1: **Input:** coarse-grained configuration $\mathbf{R}$, velocity field $\mathbf{v}^\theta$, noise distribution $\pi_{\epsilon|R}$
- 2: Sample noise: $\epsilon \sim \pi_{\epsilon|R}$
- 3: Define: $\mathbf{x}_0 = [\mathbf{R} \quad \epsilon]^\top$
- 4: Numerically solve Equation 12:

$$\mathbf{x}_1 = \mathbf{x}_0 + \int_0^1 d\tau \mathbf{v}_\tau^\theta(\phi_\tau(\mathbf{x}_0))$$

- 5: **Output:** Sample on fiber $\mathbf{r} = \mathbf{x}_1 \in \mathcal{F}(\mathbf{R})$
-

As outlined in Appendix A.2, we can write the fiber average of an observable $O : \mathbb{R}^n \rightarrow \mathbb{R}^d$ using Equation 8, as

$$\mathbb{E}_{\mathbf{r}|\mathbf{R}}[O(\mathbf{r})] = \mathbb{E}_{\epsilon|R}[O(\phi_1(\mathbf{R}, \epsilon))]. \quad (18)$$

By combining Equations 16 and 18, we can use split-flows to obtain an estimate of the local mapping entropy in Equation 9:

$$S(\mathbf{R}) = -k_B \mathbb{E}_{\epsilon|R} [\log \pi_{\epsilon|R}(\epsilon | \mathbf{R})] + k_B \mathbb{E}_{\epsilon|R} \left[\int_0^1 d\tau \nabla \cdot \mathbf{v}_\tau^\theta(\phi_\tau(\mathbf{R}, \epsilon)) \right], \quad (19)$$

which requires evaluating the entropy of the noise distribution and the volume change under the flow.

5 EXPERIMENTS

In the experimental section, we present split-flows across three molecular systems. First, we consider the mini-protein chignolin, where we validate its backmapping capabilities and quantify the information loss along a coarse-grained molecular dynamics (MD) trajectory. Next, we investigate information loss in a coarse-grained representation of a two-particle solute dragged through a lipid membrane. Finally, we present the information loss landscape in the Ramachandran representation of alanine dipeptide.

5.1 Chignolin

We apply split-flows to chignolin, a protein composed of 10 amino acids and 77 heavy atoms, which, despite its manageable size, already exhibits folding behavior. The coarse-graining map reduces the fine-grained configuration to the 10 C_α atoms, a commonly used reduction for proteins, as depicted in Figure 3.

We train split-flows on 50k frames from a 1 μ s atomistic MD simulation at 360 K, which includes multiple folding and unfolding transitions. Since split-flows

operate on Cartesian coordinates, we use the $E(3)$ -equivariant graph neural network (GNN) architecture proposed by [Satorras et al., 2021]. As the noise distribution $\pi_{\epsilon|R}$, we choose a residue-wise Gaussian distribution centered at the position of the respective C_α atom. A detailed description and the hyperparameters for both the simulation and the model are provided in Appendix B.1.

Backmapping: First, we validate split-flows by means of backmapping. Given a test set of atomistic and coarse-grained configurations, $\mathbf{r}$ and $\mathbf{R} = M(\mathbf{r})$, we sample reconstructions $\hat{\mathbf{r}} = \phi_1(\mathbf{R}, \epsilon)$ with $\epsilon \sim \pi_{\epsilon|R}$. We compare split-flows to three methods: TC-VAE [Shmilovich et al., 2022], Flow-back [Jones et al., 2025], and CG-back [Torre and Sugita, 2025]. Flow-back and CG-back transfer to unseen molecules via residue-based backmapping of the C_α -representation, so we use the authors’ pretrained models. We retrain TC-VAE with the released code and hyperparameters.

To evaluate structural fidelity, we project the fine-grained configurations onto the first two components of a time-lagged independent component analysis (TICA) [Pérez-Hernández et al., 2013], a commonly used projection. In Figure 4, we compare the resulting log-densities in this two-dimensional representation. We find that split-flows, aside from slight smoothing, reproduce all major modes of the original density—especially the misfolded state, which is typically underrepresented by other methods—indicating high diversity of backmapped samples.

In Table 1, we report several numerical metrics. We measure energetic plausibility by computing the Wasserstein-1 distance W_1 between the distributions of internal energies of the original and reconstructed configurations. We assess the consistency of backmapped configurations with their initial coarse-grained counterparts by calculating the root-mean-squared deviation (RMSD) between the original coarse-grained representation $\mathbf{R} = M(\mathbf{r})$ and the projected backmapped configuration $\hat{\mathbf{R}} = M(\hat{\mathbf{r}})$, denoted by RMSD_{cg} . To evaluate topological agreement, we construct a molecular graph based on atomic distances and compute the relative graph edit distance D_G with respect to the true molecular graph. For these three metrics we report mean and standard deviation over five test trajectories, each containing 10k frames.

To measure diversity within a fiber, we generate a set of configurations $\hat{\mathbf{r}}_i \in \mathcal{F}(M(\mathbf{r}))$ for a given reference structure $\mathbf{r}$, and compute the average RMSD between generated configurations and the reference, denoted as RMSD_{ref} , as well as the average pairwise RMSD between the generated configurations, denoted

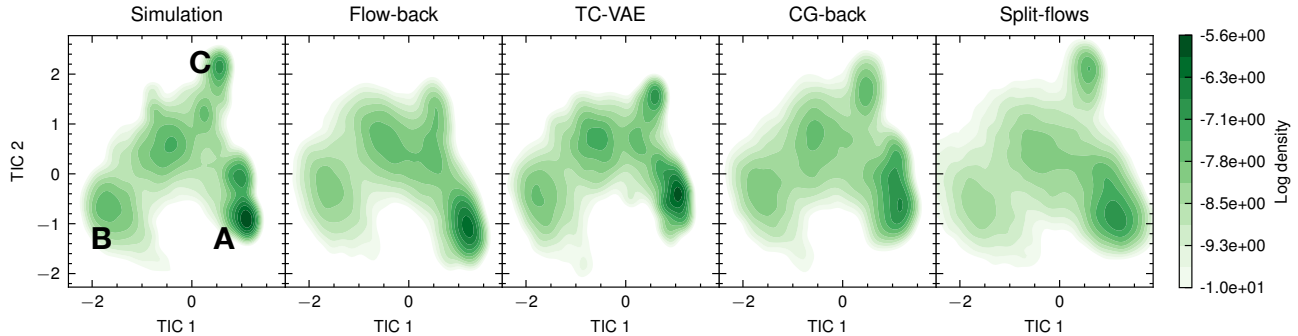


Figure 4: Log densities in the plane of the first two components of TICA. We present the projected log densities of the original simulated configurations as well as backmapped configurations using reference methods and our split-flows. The projection separates the folded (A), unfolded (B), and misfolded (C) modes of chignolin.

as RMSD_{gen} . Following [Jones et al., 2023], we define a diversity score η_{div} as the ratio $\text{RMSD}_{\text{gen}}/\text{RMSD}_{\text{ref}}$. This ratio vanishes for deterministic backmapping, where all generated samples are identical, and increases with sample diversity. We report mean and standard deviation over 50 reference configurations, with 1k samples on the fiber per reference.

Table 1: We report the Wasserstein-1 distance W_1 of the internal energy distribution in kcal mol^{-1} , the RMSD in the coarse-grained space RMSD_{cg} in pm, the relative graph edit distance D_G in %, and the fiber-diversity score η_{div} . Best values are highlighted in bold, second-best values are underlined.

Model	$W_1(\downarrow)$	$\text{RMSD}_{\text{cg}}(\downarrow)$	$D_G(\downarrow)$	$\eta_{\text{div}}(\uparrow)$
Flow-back	300 ± 3	4.602 ± 0.008	0.027 ± 0.006	0.60 ± 0.10
TC-VAE	5900 ± 150	5.0 ± 0.3	6.2 ± 0.9	0.022 ± 0.005
CG-back	321 ± 3	0.071 ± 0.007	0.53 ± 0.07	0.90 ± 0.09
Split-flows	131.1 ± 1.8	<u>0.62 ± 0.04</u>	<u>0.22 ± 0.06</u>	<u>0.79 ± 0.15</u>

Across all numerical metrics in Table 1, we find that split-flows perform competitively compared to existing methods. In particular, their ability to compute highly diverse samples—with a diversity score of 0.79—that are simultaneously energetically plausible, with a Wasserstein-1 distance of $131.1 \text{ kcal mol}^{-1}$, places split-flows in a prominent position in the comparison. Moreover, split-flows rank second in terms of coarse-grained consistency and relative graph edit distance. Nonetheless, we emphasize that Flow-back and CG-back exhibit consistently strong performance, despite their transferability. We note that the low diversity score for TC-VAE results from the model’s coherency with the previous atomistic configuration, which limits generated configurations to remain temporally consistent with their respective predecessors. Furthermore, despite granting it a multiple of the training time used for our method, we find that TC-VAE does not perform as well as presented in the orig-

inal work. More details on training and the compute budget used can be found in Appendix B.1.

Information loss: Next, we leverage the mapping entropy framework developed in Sections 3.3 and 4 to quantify the local information loss of the coarse-grained representation along a MD trajectory. We compute the information loss over a short section of a test trajectory and visualize the resulting sequence in Figure 5.

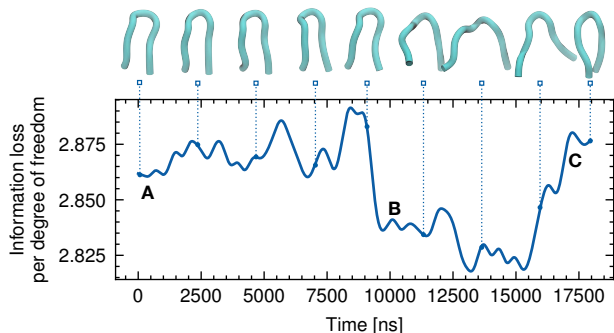


Figure 5: Average information loss per removed degree of freedom in the C_α representation of chignolin along a MD trajectory. We analyze a short section of the simulation starting in a folded state (A), followed by a partial separation of the two strands (B), and returning to the folded state (C).

The reduction to the C_α atoms, as depicted in Figure 3, projects out many orthogonal degrees of freedom, particularly in the side chains. The associated removal of interactions leads to a conformation-dependent information loss: We observe a drop in the information loss landscape in the region where the two strands of the protein separate. This partial opening reduces the interactions between the projected-out atoms in the tails, resulting in a less constrained—and therefore less informative—fiber distribution.

5.2 Solute in a Lipid Bilayer

Since the solute is approximately a rigid body, its configuration is well-described by a two-dimensional description consisting of the distance $z \in [-\frac{L}{2}, \frac{L}{2}]$ of the solute’s center of mass from the membrane center and the relative orientation $\vartheta \in [0, \pi]$ with respect to the z -axis, as depicted in Figure 6. We define a coarse-grained description of the solute by projecting out the rotational degree of freedom: $M : [-\frac{L}{2}, \frac{L}{2}] \times [0, \pi] \rightarrow [-\frac{L}{2}, \frac{L}{2}]$. We then train a split-flow, parameterized by a multilayer perceptron (MLP), to connect the configurations $\mathbf{r} = [z \ \vartheta]^\top$ and $\mathbf{R} = [z]^\top$ and use a uniform distribution on $[0, \pi]$, $\pi_{\epsilon|R} = u_{[0, \pi]}$, for the noise dimensions. Periodicity is enforced by a simple sine-cosine input parameterization for the MLP.

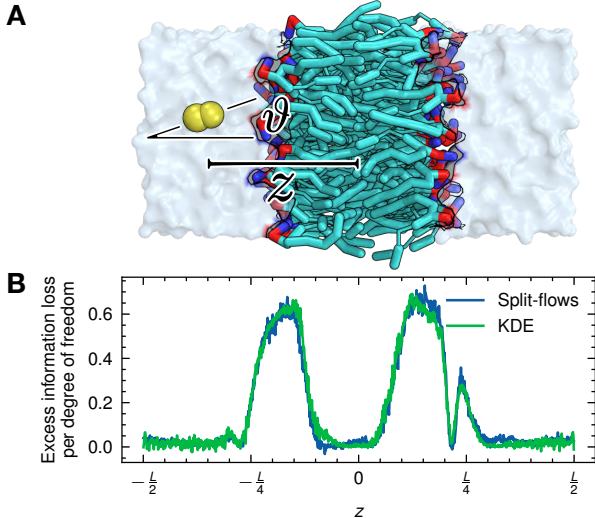


Figure 6: (A) An amphiphilic solute is dragged through a lipid bilayer surrounded by bulk water under a constant driving force. We describe its configuration by the distance z to the membrane center and its relative orientation ϑ with respect to the z -axis. (B) Average excess information loss per degree of freedom when removing ϑ , shown as a function of z .

As a baseline, we estimate the fine- and coarse-grained densities using a simple kernel density estimator (KDE), yielding a binned approximation of the local information loss; see Appendix B.2 for details. Figure 6 shows the resulting excess information loss as a function of z . The split-flow estimates closely match the KDE baseline in both shape (Pearson correlation 0.99) and local magnitude (mean absolute error 0.027) across the coarse-grained domain.

The landscape reflects the amphiphilic interactions between the lipid membrane and the solute, and the associated constraints on the solute’s relative orientation. In bulk water, these interactions are weak, and the so-

lute’s orientation is largely unconstrained, resulting in vanishing information loss. Near the surface, the hydrophilic headgroups attract the hydrophilic and repel the hydrophobic side of the solute, aligning it with the surface normal and causing a small peak in the information loss. Upon entering the membrane, the solute flips and orients its hydrophobic side toward the membrane interior. This re-orientation strongly constrains the solute, leading to a pronounced maximum in information loss at the interface. In the hydrophobic core, the constraint relaxes as both orientations become nearly equivalent, resulting in a clear decrease in information loss toward the bilayer midplane. This behavior is conceptually mirrored across the midplane, with a quantitative asymmetry due to the solute being pulled through the membrane with constant force.

5.3 Alanine Dipeptide

We consider 50k frames from a 1 μ s MD simulation of alanine dipeptide at 600K. We train the $E(3)$ -equivariant GNN parameterization of split-flows, introduced in Section 5.1, to connect the atomistic configuration with a reduced description in which only the five backbone atoms defining the dihedral angles (ϕ, ψ) —the Ramachandran angles—are retained. This coarse-graining scheme is illustrated in Figure 7. We provide a detailed description of the simulation and model hyperparameters in Appendix B.3.

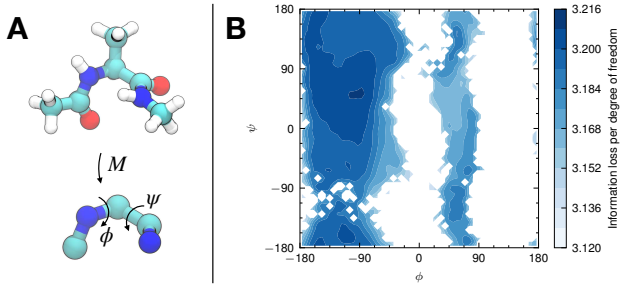


Figure 7: (A) The coarse-graining map reduces the atomistic configuration to the five atoms defining the Ramachandran dihedrals. (B) A landscape of the average information loss per removed degree of freedom is shown in the (ϕ, ψ) -plane.

Because the bond lengths and angles within this fragment are nearly rigid, the dihedrals (ϕ, ψ) uniquely determine the coarse-grained configuration $\mathbf{R}$, up to global rigid-body motions corresponding to the Euclidean group $E(3)$. This enables us to visualize the information loss landscape over the two-dimensional (ϕ, ψ) plane in Figure 7. This coarse-grained representation produces a complex distribution of lost information across the Ramachandran plane, reflecting the interactions of the eliminated degrees of freedom,

including steric repulsions that generate the forbidden regions (white) and dipole–dipole interactions that shape the overall conformational preferences of the dipeptide. It demonstrates our model’s ability to resolve highly non-trivial structure in the information loss of coarse-grained representations.

6 CONCLUSION

In this paper, we present split-flows, a novel approach for connecting molecular densities at different resolutions. Split-flows perform competitively in backmapping and—leveraging the volume change under the flow—can quantify the local information loss across general biophysical systems.

Our method performs well on various molecular systems. To scale up the method to larger macromolecules, autoregressive techniques—already used in the context of residue-based backmapping methods—may offer an appealing strategy. We propose to explore this direction in future work, where our contribution would provide unique insight in the scaling of resolution-based information loss.

Applications of our method are widespread. In particular, statistical thermodynamics provides a rich set of physical quantities linked to local mapping entropy, including the specific heat of a coarse configuration [Foley et al., 2015]. In addition, split-flows offer a principled approach for identifying informative coarse-grained representations of molecular systems—specifically, those with low and uniform information loss. Finally, using split-flows to construct scale transitions in multi-scale molecular simulations represents a promising direction for future work.

Acknowledgements

We thank Daniel Nagel for providing the lipid bilayer simulation and for fruitful discussions. Furthermore, we gratefully acknowledge William Noid for providing extensive and constructive feedback on an earlier preprint of this work. This work is supported by Deutsche Forschungsgemeinschaft (DFG) under Germany’s Excellence Strategy EXC-2181/1–390900948 (the Heidelberg STRUCTURES Excellence Cluster). The authors acknowledge support by the state of Baden-Württemberg through bwHPC and the German Research Foundation (DFG) through grant INST 35/1597-1 FUGG.

References

[Albergo et al., 2023] Albergo, M. S., Boffi, N. M., and Vanden-Eijnden, E. (2023). Stochastic inter-

polants: A unifying framework for flows and diffusions. *arXiv preprint arXiv:2303.08797*.

[Albergo et al., 2024] Albergo, M. S., Goldstein, M., Boffi, N. M., Ranganath, R., and Vanden-Eijnden, E. (2024). Stochastic interpolants with data-dependent couplings. In Salakhutdinov, R., Kolter, Z., Heller, K., Weller, A., Oliver, N., Scarlett, J., and Berkenkamp, F., editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 921–937. PMLR.

[Armstrong et al., 2012] Armstrong, J. A., Chakravarty, C., and Ballone, P. (2012). Statistical mechanics of coarse graining: Estimating dynamical speedups from excess entropies. *Journal of Chemical Physics*, 136.

[Berlaga et al., 2025] Berlaga, A., Jones, M. S., and Ferguson, A. L. (2025). Flowback-adjoint: Physics-aware and energy-guided conditional flow-matching for all-atom protein backmapping. *arXiv preprint arXiv:2508.03619*.

[Brehmer and Cranmer, 2020] Brehmer, J. and Cranmer, K. (2020). Flows for simultaneous manifold learning and density estimation. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H., editors, *Advances in Neural Information Processing Systems*, volume 33, pages 442–453. Curran Associates, Inc.

[Chen et al., 2018] Chen, R. T. Q., Rubanova, Y., Bettencourt, J., and Duvenaud, D. K. (2018). Neural ordinary differential equations. In Bengio, S., Wallach, H., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R., editors, *Advances in Neural Information Processing Systems*, volume 31. Curran Associates, Inc.

[Foley et al., 2015] Foley, T. T., Shell, M. S., and Noid, W. G. (2015). The impact of resolution upon entropy and information in coarse-grained models. *The Journal of Chemical Physics*, 143.

[Giulini et al., 2024] Giulini, M., Fiorentini, R., Tubiana, L., Potestio, R., and Menichetti, R. (2024). Excogito, an extensible coarse-graining toolbox for the investigation of biomolecules by means of low-resolution representations. *Journal of Chemical Information and Modeling*, 64:4912–4927.

[Giulini et al., 2020] Giulini, M., Menichetti, R., Shell, M. S., and Potestio, R. (2020). An information-theory-based approach for optimal model reduction of biomolecules. *Journal of Chemical Theory and Computation*, 16.

- [Holtzman et al., 2022] Holtzman, R., Giulini, M., and Potestio, R. (2022). Making sense of complex systems through resolution, relevance, and mapping entropy. *Physical Review E*, 106:044101.
- [Jin et al., 2023] Jin, J., Schweizer, K. S., and Voth, G. A. (2023). Understanding dynamics in coarse-grained models. i. universal excess entropy scaling relationship. *Journal of Chemical Physics*, 158.
- [Jones et al., 2025] Jones, M. S., Khanna, S., and Ferguson, A. L. (2025). Flowback: A generalized flow-matching approach for biomolecular backmapping. *Journal of Chemical Information and Modeling*, 65:672–692.
- [Jones et al., 2023] Jones, M. S., Shmilovich, K., and Ferguson, A. L. (2023). Diamondback: Diffusion-denoising autoregressive model for non-deterministic backmapping of α protein traces. *Journal of Chemical Theory and Computation*, 19.
- [Kidder and Noid, 2024] Kidder, K. M. and Noid, W. G. (2024). Analysis of mapping atomic models to coarse-grained resolution. *The Journal of Chemical Physics*, 161.
- [Kidder et al., 2021] Kidder, K. M., Szukalo, R. J., and Noid, W. G. (2021). Energetic and entropic considerations for coarse-graining. *The European Physical Journal B*, 94:153.
- [Lee, 2003] Lee, J. M. (2003). *Smooth Manifolds*, pages 1–29. Springer New York, New York, NY.
- [Li et al., 2020] Li, W., Burkhart, C., Polńska, P., Harmandaris, V., and Doxastakis, M. (2020). Backmapping coarse-grained macromolecules: An efficient and versatile machine learning approach. *The Journal of Chemical Physics*, 153.
- [Lipman et al., 2023] Lipman, Y., Chen, R. T. Q., Ben-Hamu, H., Nickel, M., and Le, M. (2023). Flow matching for generative modeling. In *The Eleventh International Conference on Learning Representations*.
- [Mussi and Noid, 2025] Mussi, L. M. and Noid, W. G. (2025). Predicting energetic and entropic driving forces with coarse-grained models. *The Journal of Chemical Physics*, 163.
- [Nagel and Bereau, 2025] Nagel, D. and Bereau, T. (2025). Fokker-planck score learning: Efficient free-energy estimation under periodic boundary conditions. *arXiv preprint arXiv:2506.15653*.
- [Noid, 2013] Noid, W. G. (2013). Perspective: Coarse-grained models for biomolecular systems. *The Journal of Chemical Physics*, 139.
- [Noid, 2023] Noid, W. G. (2023). Perspective: Advances, challenges, and insight for predictive coarse-grained models. *Journal of Physical Chemistry B*, 127:4174–4207.
- [Peter and Kremer, 2009] Peter, C. and Kremer, K. (2009). Multiscale simulation of soft matter systems - from the atomistic to the coarse-grained level and back. *Soft Matter*, 5.
- [Pimentel and Mandelshtam, 2025] Pimentel, J. V. and Mandelshtam, V. A. (2025). From all-atom to rigid monomer treatment of molecular clusters. *Journal of Chemical Physics*, 162.
- [Pérez-Hernández et al., 2013] Pérez-Hernández, G., Paul, F., Giorgino, T., Fabritiis, G. D., and Noé, F. (2013). Identification of slow molecular order parameters for markov model construction. *Journal of Chemical Physics*, 139.
- [Rezende and Mohamed, 2015] Rezende, D. J. and Mohamed, S. (2015). Variational inference with normalizing flows. In *32nd International Conference on Machine Learning, ICML 2015*, volume 2.
- [Rudzinski and Noid, 2011] Rudzinski, J. F. and Noid, W. G. (2011). Coarse-graining entropy, forces, and structures. *Journal of Chemical Physics*, 135.
- [Rzepiela et al., 2010] Rzepiela, A. J., Schäfer, L. V., Goga, N., Risselada, H. J., Vries, A. H. D., and Marrink, S. J. (2010). Software news and update reconstruction of atomistic details from coarse-grained structures. *Journal of Computational Chemistry*, 31.
- [Satorras et al., 2021] Satorras, V. G., Hoogeboom, E., and Welling, M. (2021). E(n) equivariant graph neural networks. In Meila, M. and Zhang, T., editors, *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pages 9323–9332. PMLR.
- [Shell, 2008] Shell, M. S. (2008). The relative entropy is fundamental to multiscale and inverse thermodynamic problems. *The Journal of Chemical Physics*, 129.
- [Shmilovich et al., 2022] Shmilovich, K., Stieffenhofer, M., Charron, N. E., and Hoffmann, M. (2022). Temporally coherent backmapping of molecular trajectories from coarse-grained to atomistic resolution. *The Journal of Physical Chemistry A*, 126:9124–9139.
- [Stieffenhofer et al., 2020] Stieffenhofer, M., Wand, M., and Bereau, T. (2020). Adversarial reverse

mapping of equilibrated condensed-phase molecular structures. *Machine Learning: Science and Technology*, 1:045014.

- [Tong et al., 2024] Tong, A. Y., Malkin, N., Fattas, K., Atanackovic, L., Zhang, Y., Huguet, G., Wolf, G., and Bengio, Y. (2024). Simulation-free Schrödinger bridges via score and flow matching. In Dasgupta, S., Mandt, S., and Li, Y., editors, *Proceedings of The 27th International Conference on Artificial Intelligence and Statistics*, volume 238 of *Proceedings of Machine Learning Research*, pages 1279–1287. PMLR.
- [Torre and Sugita, 2025] Torre, D. U. L. and Sugita, Y. (2025). Cgback: Diffusion model for backmapping large-scale and complex coarse-grained molecular systems. *Journal of Chemical Information and Modeling*.
- [Wang et al., 2022] Wang, W., Xu, M., Cai, C., Miller, B. K., Smidt, T., Wang, Y., Tang, J., and Gomez-Bombarelli, R. (2022). Generative coarse-graining of molecular conformations. In Chaudhuri, K., Jegelka, S., Song, L., Szepesvari, C., Niu, G., and Sabato, S., editors, *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, pages 23213–23236. PMLR.
- [Wassenaar et al., 2014] Wassenaar, T. A., Pluhackova, K., Böckmann, R. A., Marrink, S. J., and Tieleman, D. P. (2014). Going backward: A flexible geometric approach to reverse transformation from coarse grained to atomistic models. *Journal of Chemical Theory and Computation*, 10.

Checklist

1. For all models and algorithms presented, check if you include:
 - (a) A clear description of the mathematical setting, assumptions, algorithm, and/or model. [Yes]
 - (b) An analysis of the properties and complexity (time, space, sample size) of any algorithm. [Yes, see related work for details]
 - (c) (Optional) Anonymized source code, with specification of all dependencies, including external libraries. [Yes, available at <https://github.com/BereauLab/split-flows>]
2. For any theoretical claim, check if you include:
 - (a) Statements of the full set of assumptions of all theoretical results. [Yes]
 - (b) Complete proofs of all theoretical results. [Yes]
 - (c) Clear explanations of any assumptions. [Yes]
3. For all figures and tables that present empirical results, check if you include:
 - (a) The code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL). [Yes, code is available at <https://github.com/BereauLab/split-flows>, data is available upon request.]
 - (b) All the training details (e.g., data splits, hyperparameters, how they were chosen). [Yes, see corresponding sections of the Appendix B.]
 - (c) A clear definition of the specific measure or statistics and error bars (e.g., with respect to the random seed after running experiments multiple times). [Yes]
 - (d) A description of the computing infrastructure used. (e.g., type of GPUs, internal cluster, or cloud provider). [Yes, see corresponding sections of the Appendix B.]
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets, check if you include:
 - (a) Citations of the creator If your work uses existing assets. [Yes]
 - (b) The license information of the assets, if applicable. [Yes]
 - (c) New assets either in the supplemental material or as a URL, if applicable. [Yes, data is available upon request.]
 - (d) Information about consent from data providers/curators. [Yes]
 - (e) Discussion of sensible content if applicable, e.g., personally identifiable information or offensive content. [Not Applicable]
5. If you used crowdsourcing or conducted research with human subjects, check if you include:
 - (a) The full text of instructions given to participants and screenshots. [Not Applicable]
 - (b) Descriptions of potential participant risks, with links to Institutional Review Board (IRB) approvals if applicable. [Not Applicable]
 - (c) The estimated hourly wage paid to participants and the total amount spent on participant compensation. [Not Applicable]

Split-Flows: Measure Transport and Information Loss Across Molecular Resolutions: Supplementary Materials

A ADDITIONAL PROOFS AND THEORETICAL DETAILS

In this section, we provide additional theoretical details that complement the derivations of the local mapping entropy, information loss, and the split-flow setup in the main text. The propositions and proofs presented here rely heavily on well-established theory on normalizing flows [Rezende and Mohamed, 2015, Chen et al., 2018, Lipman et al., 2023, Albergo et al., 2023]. For completeness, we briefly state the core results.

Let $\phi_1 : \mathbb{R}^n \rightarrow \mathbb{R}^n$ be a diffeomorphism and π_0 a probability density on $\mathbb{R}^n$. The density of the pushforward measure $(\phi_1)_\# \pi_0$ under the flow then can be written as

$$(\phi_1)_\# \pi_0(\mathbf{x}_0) = |\det J_{\phi_1}(\mathbf{x}_0)|^{-1} \pi_0(\mathbf{x}_0) = \pi_1(\phi_1(\mathbf{x}_0)), \quad (20)$$

where $J_{\phi_1}(\mathbf{x}_0)$ is the Jacobian matrix of partial derivatives. In case the flow ϕ_1 is parameterized in continuous time by the ordinary differential equation (ODE):

$$\frac{d}{dt} \phi_t(\mathbf{x}_0) = \mathbf{v}_t(\phi_t(\mathbf{x}_0)), \quad \phi_0(\mathbf{x}_0) = \mathbf{x}_0, \quad (21)$$

then the logarithm of the Jacobian determinant evolves according to the ODE:

$$\frac{d}{dt} \log |\det J_{\phi_t}(\mathbf{x}_0)| = \nabla \cdot \mathbf{v}_t(\phi_t(\mathbf{x}_0)), \quad (22)$$

which yields an integral expression for the total change of volume along the flow:

$$\log |\det J_{\phi_1}(\mathbf{x}_0)| = \int_0^1 d\tau \nabla \cdot \mathbf{v}_\tau(\phi_\tau(\mathbf{x}_0)). \quad (23)$$

Consequently, the pushforward density can be expressed as

$$\pi_1(\phi_1(\mathbf{x}_0)) = \pi_0(\mathbf{x}_0) \exp \left[- \int_0^1 d\tau \nabla \cdot \mathbf{v}_\tau(\phi_\tau(\mathbf{x}_0)) \right]. \quad (24)$$

Together, these results formalize how probability densities evolve under smooth transformations and serve as the starting point for our theoretical analysis of mapping entropy, information loss, and split-flows.

A.1 Decomposition of the Coarse-Grained Potential

We now show that the coarse-grained potential of mean force admits a natural decomposition into energetic and entropic contributions.

Proposition A.1 (Decomposition of the coarse-grained potential). *Let $\mathbf{r} \in \mathbb{R}^n$ denote a fine-grained configuration, and let $\mathbf{R} = M(\mathbf{r}) \in \mathbb{R}^N$ be the associated coarse-grained representative obtained by a measurable coarse-graining map $M : \mathbb{R}^n \rightarrow \mathbb{R}^N$. Suppose the fine-grained configurations are Boltzmann-distributed as*

$$\pi_r(\mathbf{r}) = Z^{-1} \exp[-u(\mathbf{r})/(k_B T)], \quad (25)$$

where $u(\mathbf{r})$ is the potential energy governing the fine-grained distribution and Z is the partition function. Then the marginal coarse-grained distribution $\pi_R(\mathbf{R})$ defined by

$$\pi_R(\mathbf{R}) = \int_{\mathbb{R}^n} d\mathbf{r} \delta(\mathbf{R} - M(\mathbf{r})) \pi_r(\mathbf{r}) \quad (26)$$

can be written in Boltzmann form:

$$\pi_R(\mathbf{R}) \propto \exp[-W(\mathbf{R})/(k_B T)], \quad (27)$$

where the free energy $W(\mathbf{R})$ —the potential of mean force (PMF)—admits the decomposition

$$W(\mathbf{R}) = E(\mathbf{R}) - TS(\mathbf{R}), \quad (28)$$

with

$$E(\mathbf{R}) = \mathbb{E}_{r|R}[u(\mathbf{r})], \quad S(\mathbf{R}) = -k_B \mathbb{E}_{r|R} \left[\log \frac{\pi_r(\mathbf{r})}{\pi_R(\mathbf{R})} \right]. \quad (29)$$

Proof. Starting from the Boltzmann form of the coarse-grained density,

$$\pi_R(\mathbf{R}) \propto \exp[-W(\mathbf{R})/(k_B T)], \quad (30)$$

we express the PMF as

$$W(\mathbf{R}) = -k_B T \log \pi_R(\mathbf{R}) + \text{const.} \quad (31)$$

Similarly, from the fine-grained Boltzmann distribution we can write

$$u(\mathbf{r}) = -k_B T \log \pi_r(\mathbf{r}) + \text{const.} \quad (32)$$

Taking the conditional expectation of equation 32 over the conditional distribution $\pi_{r|R}$, we obtain

$$\mathbb{E}_{r|R}[u(\mathbf{r})] = -k_B T \mathbb{E}_{r|R}[\log \pi_r(\mathbf{r})] + \text{const.} \quad (33)$$

Inserting Equation 33 into the terms of Equation 28, written in Equation 29, we find:

$$W(\mathbf{R}) = E(\mathbf{R}) - TS(\mathbf{R}) \quad (34)$$

$$= \mathbb{E}_{r|R}[u(\mathbf{r})] + k_B T \mathbb{E}_{r|R} \left[\log \frac{\pi_r(\mathbf{r})}{\pi_R(\mathbf{R})} \right] \quad (35)$$

$$= \mathbb{E}_{r|R}[u(\mathbf{r})] - \mathbb{E}_{r|R}[u(\mathbf{r})] - k_B T \mathbb{E}_{r|R}[\log \pi_R(\mathbf{R})] + \text{const.} \quad (36)$$

$$= -k_B T \mathbb{E}_{r|R}[\log \pi_R(\mathbf{R})] + \text{const.}, \quad (37)$$

which recovers Equation 31 up to the unknown constant offset. However, since potentials are defined only up to an additive constant, we can drop the constant, completing the proof. $\square$

Remark A.1.1. The decomposition $W = E - TS$ expresses the PMF as the conditional free energy of the fine-grained system constrained to the coarse-grained configuration $\mathbf{R}$. Here, $E(\mathbf{R})$ denotes the mean internal energy of the fine-grained microstates compatible with $\mathbf{R}$, while $S(\mathbf{R})$ quantifies their configurational entropy.

Remark A.1.2. The entropic contribution to the PMF, $S(\mathbf{R})$, can be identified with the local mapping entropy $S(\mathbf{R})$, which quantifies the information loss in a coarse-grained representation. Higher information loss corresponds to a lower mapping entropy and thus a higher value of the PMF, which in turn lowers the probability of the coarse-grained configuration.

A.2 Computation of Fiber Averages with Split-Flows

Split flows allow us to directly access fiber averages, i.e., expectation values of observables defined on the fine-grained space $\mathbb{R}^n$, restricted to the fiber $\mathcal{F}(\mathbf{R})$ associated with a coarse-grained representative $\mathbf{R}$. We formalize the expression stated in Equation 18 in the main text with the following proposition.

Proposition A.2 (Computation of fiber averages with split-flows). *Let π_R be a probability density on $\mathbb{R}^N$ and $\pi_{\epsilon|R}$ a conditional density on $\mathbb{R}^{n-N}$, defining a joint density $\pi_R \times \pi_{\epsilon|R}$ on $\mathbb{R}^n = \mathbb{R}^N \times \mathbb{R}^{n-N}$. Let $M : \mathbb{R}^n \rightarrow \mathbb{R}^N$ be a measurable coarse-graining map, and let $\phi_1 : \mathbb{R}^n \rightarrow \mathbb{R}^n$ be a diffeomorphism satisfying*

$$(\phi_1)_\# \pi_R(\mathbf{R}) \pi_{\epsilon|R}(\boldsymbol{\epsilon}) = |\det J_{\phi_1}(\mathbf{R}, \boldsymbol{\epsilon})|^{-1} \pi_R(\mathbf{R}) \pi_{\epsilon|R}(\boldsymbol{\epsilon} | \mathbf{R}) = \pi_r(\phi_1(\mathbf{R}, \boldsymbol{\epsilon})), \quad (38)$$

where π_r is a target density on $\mathbb{R}^n$. Assume further that ϕ_1 inverts the coarse-graining map in the sense that

$$M \circ \phi_1(\mathbf{R}, \epsilon) = \mathbf{R}, \quad (39)$$

for all $(\mathbf{R}, \epsilon) \in \mathbb{R}^N \times \mathbb{R}^{n-N}$. Finally, let $O : \mathbb{R}^n \rightarrow \mathbb{R}^d$ be a measurable observable. We can then write the conditional expectation of O over $\pi_{r|R}$ as:

$$\mathbb{E}_{r|R}[O(\mathbf{r})] = \mathbb{E}_{\epsilon|R}[O(\phi_1(\mathbf{R}, \epsilon))]. \quad (40)$$

Proof. Let π_r , π_R , $\pi_{\epsilon|R}$, ϕ_1 , M and O obey the properties above. By definition, we can write the fiber average of O for a given coarse-grained representative $\mathbf{R}$ as

$$\mathbb{E}_{r|R}[O(\mathbf{r})] = \frac{\int_{\mathbb{R}^n} d\mathbf{r} \delta(M(\mathbf{r}) - \mathbf{R}) O(\mathbf{r}) \pi_r(\mathbf{r})}{\int_{\mathbb{R}^n} d\mathbf{r} \delta(M(\mathbf{r}) - \mathbf{R}) \pi_r(\mathbf{r})}. \quad (41)$$

By substituting $\mathbf{r} = \phi_1(\mathbf{R}', \epsilon)$ and using the change-of-variables theorem $d\mathbf{r} = |\det J_{\phi_1}(\mathbf{R}', \epsilon)| d\mathbf{R}' d\epsilon$, we can rewrite the expectation as

$$\mathbb{E}_{r|R}[O(\mathbf{r})] = \frac{\int_{\mathbb{R}^n} d\mathbf{r} \delta(M(\mathbf{r}) - \mathbf{R}) O(\mathbf{r}) \pi_r(\mathbf{r})}{\int_{\mathbb{R}^n} d\mathbf{r} \delta(M(\mathbf{r}) - \mathbf{R}) \pi_r(\mathbf{r})} \quad (42)$$

$$= \frac{\int_{\mathbb{R}^N} \int_{\mathbb{R}^{n-N}} d\mathbf{R}' d\epsilon \delta(M(\phi_1(\mathbf{R}', \epsilon)) - \mathbf{R}) O(\phi_1(\mathbf{R}', \epsilon)) \pi_r(\phi_1(\mathbf{R}', \epsilon)) |\det J_{\phi_1}(\mathbf{R}', \epsilon)|}{\int_{\mathbb{R}^N} \int_{\mathbb{R}^{n-N}} d\mathbf{R}' d\epsilon \delta(M(\phi_1(\mathbf{R}', \epsilon)) - \mathbf{R}) \pi_r(\phi_1(\mathbf{R}', \epsilon)) |\det J_{\phi_1}(\mathbf{R}', \epsilon)|} \quad (43)$$

$$\stackrel{M \circ \phi_1 = \text{Id}_R}{=} \frac{\int_{\mathbb{R}^{n-N}} d\epsilon O(\phi_1(\mathbf{R}, \epsilon)) \pi_r(\phi_1(\mathbf{R}, \epsilon)) |\det J_{\phi_1}(\mathbf{R}, \epsilon)|}{\int_{\mathbb{R}^{n-N}} d\epsilon \pi_r(\phi_1(\mathbf{R}, \epsilon)) |\det J_{\phi_1}(\mathbf{R}, \epsilon)|}, \quad (44)$$

where $J_{\phi_1}(\mathbf{R}, \epsilon)$ denotes the Jacobian matrix of partial derivatives of ϕ_1 . Identifying

$$\pi_r(\phi_1(\mathbf{R}, \epsilon)) |\det J_{\phi_1}(\mathbf{R}, \epsilon)| = \pi_R(\mathbf{R}) \pi_{\epsilon|R}(\epsilon | \mathbf{R})$$

via the pushforward relation in Equation 38, we obtain

$$\mathbb{E}_{r|R}[O(\mathbf{r})] = \frac{\int_{\mathbb{R}^{n-N}} d\epsilon O(\phi_1(\mathbf{R}, \epsilon)) \pi_R(\mathbf{R}) \pi_{\epsilon|R}(\epsilon | \mathbf{R})}{\int_{\mathbb{R}^{n-N}} d\epsilon \pi_R(\mathbf{R}) \pi_{\epsilon|R}(\epsilon | \mathbf{R})} \quad (45)$$

$$= \int_{\mathbb{R}^{n-N}} d\epsilon O(\phi_1(\mathbf{R}, \epsilon)) \pi_{\epsilon|R}(\epsilon | \mathbf{R}) \quad (46)$$

$$= \mathbb{E}_{\epsilon|R}[O(\phi_1(\mathbf{R}, \epsilon))], \quad (47)$$

which completes the proof. $\square$

Remark A.2.1 (Practical estimation). In practice, we approximate the expectation value over $\pi_{\epsilon|R}$ using a Monte Carlo estimate:

$$\mathbb{E}_{\epsilon|R}[O(\phi_1(\mathbf{R}, \epsilon))] \approx \frac{1}{M} \sum_{i=1}^M O(\phi_1(\mathbf{R}, \epsilon_i)). \quad (48)$$

Here, ϵ_i are M samples drawn from the noise distribution $\pi_{\epsilon|R}$, which is straightforward to sample from by construction.

A.3 Mapping Entropy Estimation with Split-Flows

Split-flows allow us to obtain an unbiased estimate of the local mapping entropy $S(\mathbf{R})$ for a given coarse-grained configuration $\mathbf{R}$. This section will complement the derivations of Equation 19 in the main text with a formal proposition, showing that the estimator arises naturally from the pushforward structure of the flow.

Proposition A.3 (Mapping entropy estimation with split-flows). *Let π_R be a probability density on $\mathbb{R}^N$ and $\pi_{\epsilon|R}$ a conditional density on $\mathbb{R}^{n-N}$, defining a joint density $\pi_R \times \pi_{\epsilon|R}$ on $\mathbb{R}^n = \mathbb{R}^N \times \mathbb{R}^{n-N}$. Let $M : \mathbb{R}^n \rightarrow \mathbb{R}^N$ be a measurable coarse-graining map, and let $\phi_1 : \mathbb{R}^n \rightarrow \mathbb{R}^n$ be a diffeomorphism satisfying*

$$(\phi_1)_\# \pi_R(\mathbf{R}) \pi_{\epsilon|R}(\epsilon) = |\det J_{\phi_1}(\mathbf{R}, \epsilon)|^{-1} \pi_R(\mathbf{R}) \pi_{\epsilon|R}(\epsilon | \mathbf{R}) = \pi_r(\phi_1(\mathbf{R}, \epsilon)), \quad (49)$$

where π_r is a target density on $\mathbb{R}^n$. Assume further that ϕ_1 inverts the coarse-graining map in the sense that

$$M \circ \phi_1(\mathbf{R}, \epsilon) = \mathbf{R}, \quad (50)$$

for all $(\mathbf{R}, \epsilon) \in \mathbb{R}^N \times \mathbb{R}^{n-N}$. The local mapping entropy, defined as:

$$S(\mathbf{R}) = -k_B \int_{\mathcal{F}(\mathbf{R})} d\mathbf{r} \pi_{r|R}(\mathbf{r} | \mathbf{R}) \log \pi_{r|R}(\mathbf{r} | \mathbf{R}) \quad (51)$$

then can be estimated via the split-flow setup as:

$$S(\mathbf{R}) = -k_B \mathbb{E}_{\epsilon|R} [\log \pi_{\epsilon|R}(\epsilon | \mathbf{R})] + k_B \mathbb{E}_{\epsilon|R} [\log |\det J_{\phi_1}(\mathbf{R}, \epsilon)|]. \quad (52)$$

Proof. Let ϕ_1 , π_r , and π_R fulfill the properties above. Starting from the definition of the local mapping entropy, we use Bayes' formula to rewrite the integrand:

$$S(\mathbf{R}) = -k_B \int_{\mathcal{F}(\mathbf{R})} d\mathbf{r} \pi_{r|R}(\mathbf{r} | \mathbf{R}) \log \pi_{r|R}(\mathbf{r} | \mathbf{R}) \quad (53)$$

$$= -k_B \int_{\mathcal{F}(\mathbf{R})} d\mathbf{r} \pi_{r|R}(\mathbf{r} | \mathbf{R}) \log \frac{\pi_{R|r}(\mathbf{R} | \mathbf{r}) \pi_r(\mathbf{r})}{\pi_R(\mathbf{R})} \quad (54)$$

$$= -k_B \int_{\mathcal{F}(\mathbf{R})} d\mathbf{r} \pi_{r|R}(\mathbf{r} | \mathbf{R}) \log \frac{\pi_r(\mathbf{r})}{\pi_R(M(\mathbf{r}))} \quad (55)$$

$$= -k_B \mathbb{E}_{r|R} \left[\log \frac{\pi_r(\mathbf{r})}{\pi_R(M(\mathbf{r}))} \right] \quad (56)$$

where we used that the posterior $\pi_{R|r}(\mathbf{R} | \mathbf{r}) \equiv 1$ on the integration domain $\mathcal{F}(\mathbf{R})$ and replaced $\mathbf{R}$ by $M(\mathbf{r})$. Next we are going to leverage the result of Proposition A.2 to obtain:

$$S(\mathbf{R}) = -k_B \mathbb{E}_{r|R} \left[\log \frac{\pi_r(\mathbf{r})}{\pi_R(M(\mathbf{r}))} \right] \quad (57)$$

$$= -k_B \mathbb{E}_{\epsilon|R} \left[\log \frac{\pi_r(\phi_1(\mathbf{R}, \epsilon))}{\pi_R(M(\phi_1(\mathbf{R}, \epsilon)))} \right] \quad (58)$$

$$= -k_B \mathbb{E}_{\epsilon|R} \left[\log \frac{\pi_r(\phi_1(\mathbf{R}, \epsilon))}{\pi_R(\mathbf{R})} \right]. \quad (59)$$

Inserting the pushforward relation in Equation 49 then yields:

$$S(\mathbf{R}) = -k_B \mathbb{E}_{\epsilon|R} \left[\log \frac{\pi_r(\phi_1(\mathbf{R}, \epsilon))}{\pi_R(\mathbf{R})} \right] \quad (60)$$

$$= -k_B \mathbb{E}_{\epsilon|R} [\log \pi_{\epsilon|R}(\epsilon | \mathbf{R})] + k_B \mathbb{E}_{\epsilon|R} [\log |\det J_{\phi_1}(\mathbf{R}, \epsilon)|], \quad (61)$$

which completes the proof. $\square$

Remark A.3.1 (Mapping entropy estimation with continuous normalizing flows). In case the flow ϕ_t is parameterized in continuous time by an underlying velocity field $\mathbf{v}_t$, we can replace its log Jacobian determinant with an integral over the interpolation interval $[0, 1]$ of the divergence of the flow [Lipman et al., 2023]:

$$\log |\det J_{\phi_1}(\mathbf{R}, \epsilon)| = \int_0^1 d\tau \nabla \cdot \mathbf{v}_\tau(\phi_\tau(\mathbf{R}, \epsilon)). \quad (62)$$

This recovers the estimator presented in Equation 19.

Remark A.3.2 (Variance of the mapping entropy estimator). In practice, we approximate the expectation over $\pi_{\epsilon|R}$ in the estimate of the local mapping entropy in Equation 19 using the Monte Carlo estimator presented in Remark A.2.1. In Figure 8, we analyze the variance of this estimator for the solute-in-lipid-bilayer experiment described in Section 5.2. We find that the relative variance of the estimator follows a power-law decay with respect to the number of samples.

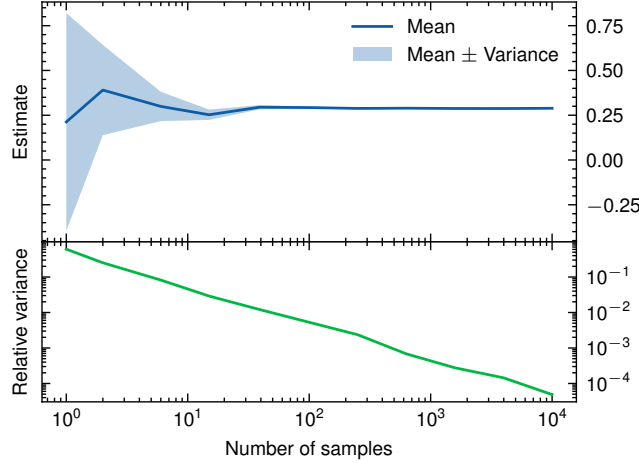


Figure 8: Estimate and variance of the local excess mapping entropy as a function of the number of samples. We show the estimate (top) and the relative variance with respect to the mean (bottom). The linear behavior of the relative variance in the log-log plot indicates a power-law decay.

A.4 Validity of the Split-Flow Coupling

The training algorithm for two-sided flow matching relies on drawing samples from the endpoint distributions. In the split-flow setup, we propose constructing a coupling based on the coarse-graining map, which encourages the flow to correctly pair fine- and coarse-grained configurations. We show that this coupling is a valid coupling.

Proposition A.4 (Validity of the split-flow coupling). *Let π_r be a probability density on $\mathbb{R}^n$. Let π_R be a probability density on $\mathbb{R}^N$ and let $\pi_{\epsilon|R}$ be a conditional density on $\mathbb{R}^{n-N}$, defining a joint density $\pi_R \times \pi_{\epsilon|R}$ on $\mathbb{R}^n = \mathbb{R}^N \times \mathbb{R}^{n-N}$. Finally, let $M : \mathbb{R}^n \rightarrow \mathbb{R}^N$ be a coarse-graining map. The joint coupling*

$$\pi_{R,\epsilon,r}(\mathbf{R}, \boldsymbol{\epsilon}, \mathbf{r}) = \pi_r(\mathbf{r}) \delta(\mathbf{R} - M(\mathbf{r})) \pi_{\epsilon|R}(\boldsymbol{\epsilon} | \mathbf{R}). \quad (63)$$

then defines a valid coupling of the two distributions π_r and $\pi_R \times \pi_{\epsilon|R}$ in the sense that the marginal over $\mathbf{r}$ is π_r , and the marginal over $(\mathbf{R}, \boldsymbol{\epsilon})$ is $\pi_R \times \pi_{\epsilon|R}$.

Proof. Let π_r , π_R , $\pi_{\epsilon|R}$, and M be defined as above. The marginal of $\pi_{R,\epsilon,r}$ over $\mathbf{r}$ then reads:

$$\int_{\mathbb{R}^n} d\mathbf{r} \int_{\mathbb{R}^{n-N}} d\boldsymbol{\epsilon} \pi_{R,\epsilon,r}(\mathbf{R}, \boldsymbol{\epsilon}, \mathbf{r}) = \int_{\mathbb{R}^N} d\mathbf{R} \int_{\mathbb{R}^{n-N}} d\boldsymbol{\epsilon} \pi_r(\mathbf{r}) \delta(\mathbf{R} - M(\mathbf{r})) \pi_{\epsilon|R}(\boldsymbol{\epsilon} | \mathbf{R}) = \pi_r(\mathbf{r}), \quad (64)$$

since the integral over $\mathbf{R}$ evaluates at $\mathbf{R} = M(\mathbf{r})$ and the inner integral over $\boldsymbol{\epsilon}$ simply integrates to 1 due to normalization of $\pi_{\epsilon|R}$. Furthermore, the marginal over $(\mathbf{R}, \boldsymbol{\epsilon})$ reads:

$$\int_{\mathbb{R}^n} d\mathbf{r} \pi_{R,\epsilon,r}(\mathbf{R}, \boldsymbol{\epsilon}, \mathbf{r}) = \underbrace{\int_{\mathbb{R}^n} d\mathbf{r} \pi_r(\mathbf{r}) \delta(\mathbf{R} - M(\mathbf{r}))}_{= \pi_R(\mathbf{R})} \pi_{\epsilon|R}(\boldsymbol{\epsilon} | \mathbf{R}) = \pi_R(\mathbf{R}) \pi_{\epsilon|R}(\boldsymbol{\epsilon} | \mathbf{R}), \quad (65)$$

where we used the definition of the coarse grained density in Equation 3. $\square$

Remark A.4.1 (Dimensionality bridging via augmentation). By augmenting the coarse-grained configuration $\mathbf{R} \in \mathbb{R}^N$ with auxiliary noise $\boldsymbol{\epsilon} \in \mathbb{R}^{n-N}$, we construct a joint variable $(\mathbf{R}, \boldsymbol{\epsilon}) \in \mathbb{R}^n$ that enables training a continuous normalizing flow $\phi_t : \mathbb{R}^n \rightarrow \mathbb{R}^n$ via flow matching. Since both endpoint distributions now are defined over the same space, the flow is well-defined and can learn a bijective transport from $\pi_R \times \pi_{\epsilon|R}$ to π_r . This approach resolves the non-invertibility of the coarse-graining map by converting the ill-posed inverse problem into a generative one.

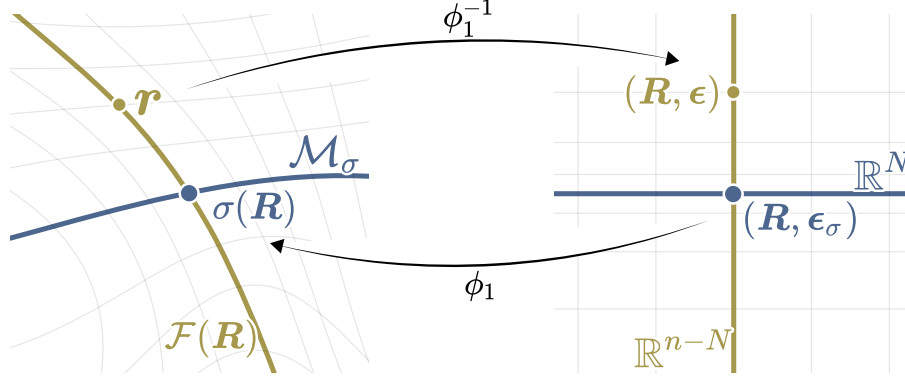


Figure 9: A geometric perspective on split-flows. The coarse-graining map M uniquely defines the fiber $\mathcal{F}(\mathbf{R})$ over a coarse-grained representation $\mathbf{R}$. By fixing a section $\sigma(\mathbf{R})$, we can define a coarse manifold $\mathcal{M}_\sigma$. Split-flows then learn a global coordinate transformation $\mathbf{r} \mapsto (\mathbf{R}, \epsilon)$, that disentangles the geometric structure induced by the coarse-graining map.

A.5 A Geometric Perspective

In this section we provide a geometric view on split-flows (see Figure 9), following ideas introduced by [Brehmer and Cranmer, 2020]. The underlying principle is to disambiguate between (i) manifold-learning, i.e., identifying the underlying manifold $\mathcal{M}_\sigma$ of slow degrees of freedom and the orthogonal fiber $\mathcal{F}(\mathbf{R})$ of fast degrees of freedom, and (ii) estimating the conditional density $\pi_{\mathbf{r}|\mathbf{R}}$ of the fast degrees of freedom on the fiber.

Fiber: We consider a smooth coarse-graining map $M : \mathbb{R}^n \rightarrow \mathbb{R}^N$, which associates fine-grained configurations $\mathbf{r} \in \mathbb{R}^n$ to a coarse-grained representative $M(\mathbf{r}) = \mathbf{R} \in \mathbb{R}^N$. The fiber $\mathcal{F}(\mathbf{R})$ over a coarse-grained representative $\mathbf{R}$ is the set

$$\mathcal{F}(\mathbf{R}) = \{\mathbf{r} \in \mathbb{R}^n \mid M(\mathbf{r}) = \mathbf{R}\}. \quad (66)$$

Its position and shape are uniquely determined by the position $\mathbf{R}$ and the mapping M . Moving along the fiber changes the fast degrees of freedom, while keeping the slow degrees of freedom fixed. Assuming that the coarse-graining map M is a smooth submersion, i.e., its Jacobian J_M has full rank, the fibers for different values of $\mathbf{R}$ form a foliation of the configurational space $\mathbb{R}^n$.

Local tangent space decomposition: Due to the submersion theorem [Lee, 2003], we can locally decompose the tangent space $T_{\mathbf{r}}\mathbb{R}^n$ into the fast (fiber) and slow (coarse) directions induced by M at $\mathbf{r} \in \mathbb{R}^n$:

$$T_{\mathbf{r}}\mathbb{R}^n = \ker J_M(\mathbf{r}) \oplus \text{range } J_M(\mathbf{r})^\top, \quad (67)$$

where $J_M(\mathbf{r}) \in \mathbb{R}^{N \times n}$ is the Jacobian matrix of partial derivatives of M , which we assume to have full rank. We can further identify the two components as local tangent spaces of the fiber $\mathcal{F}(\mathbf{R})$ and a locally defined coarse manifold $\mathcal{M}$, respectively:

$$T_{\mathbf{r}}\mathcal{F}(\mathbf{R}) = \ker J_M(\mathbf{r}), \quad T_{\mathbf{r}}\mathcal{M} = \text{range } J_M(\mathbf{r})^\top. \quad (68)$$

Gauge-freedom of the coarse manifold: While the fiber is uniquely defined, the definition of a coarse manifold $\mathcal{M}$ is ambiguous. To resolve this, we choose a section $\sigma(\mathbf{R}) \in \mathcal{F}(\mathbf{R})$ to uniquely define a coarse manifold $\mathcal{M}_\sigma$, e.g., [Pimentel and Mandelshtam, 2025] propose to use the minimum-energy configuration on the fiber. In general, the choice of section $\sigma : \mathbb{R}^N \rightarrow \mathbb{R}^n$, remains a gauge-freedom, which must be fixed to select a unique coarse manifold among the family of manifolds compatible with the local tangent-space decomposition induced by M .

Globally linearized structure: Split-flows parametrize a diffeomorphic coordinate map

$$\phi_1^{-1} : \mathbb{R}^n \rightarrow \mathbb{R}^N \times \mathbb{R}^{n-N}, \quad \mathbf{r} \mapsto (\mathbf{R}, \epsilon) \quad (69)$$

whose first N components coincide with the coarse-graining map M . In these coordinates, the fibers $\mathcal{F}(\mathbf{R})$ are mapped to affine subspaces $\{\mathbf{R}\} \times \mathbb{R}^{n-N}$, while the coarse manifold $\mathcal{M}_\sigma$ is mapped to the coordinate subspace $\mathbb{R}^N \times \{\epsilon_\sigma\}$. Thereby, split flows (i) learn a coordinate system that disentangles the geometric structure induced

by the coarse-graining map M into slow (coarse) and fast (fiber) degrees of freedom, and (ii) enable density estimation on each fiber by mapping the corresponding conditional distribution $\pi_{\mathbf{r}|R}$ on $\mathcal{F}(\mathbf{R})$ to a tractable density $\pi_{\epsilon|R}$ on the affine subspace $\{\mathbf{R}\} \times \mathbb{R}^{n-N}$.

B EXPERIMENTAL DETAILS

In this section we are going to provide additional details on the experiments performed in the experimental section of the main text. These include details on data generation, model parameterization, evaluation, and model training.

B.1 Chignolin

B.1.1 Data Generation

Langevin molecular dynamics: Molecular dynamics (MD) is a trajectory-based sampling algorithm that simulates the time evolution of the configuration $\mathbf{r}$ of a molecular system using Newton’s equations of motion. MD simulations are often performed at constant temperature T , which is maintained by coupling the system to an external heat bath—a procedure referred to as *thermostatting*. One widely used approach is the Langevin equation, which augments Newtonian dynamics with friction and stochastic thermal forces:

$$m_i \frac{d^2 \mathbf{r}_i}{dt^2} = -\nabla_{\mathbf{r}_i} u(\mathbf{r}) - \gamma m_i \frac{d\mathbf{r}_i}{dt} + \sqrt{2\gamma m_i k_B T} \boldsymbol{\zeta}_i(t), \quad (70)$$

where m_i and $\mathbf{r}_i$ denote the mass and position of the i -th particle, $u(\mathbf{r})$ is the potential energy function, γ is the friction coefficient, k_B is Boltzmann’s constant, and $\boldsymbol{\zeta}_i(t)$ is Gaussian white noise with zero mean and unit variance. The deterministic term $-\nabla_{\mathbf{r}_i} u(\mathbf{r})$ drives the system according to interatomic forces, the friction term dissipates kinetic energy, and the stochastic term restores energy from the thermal bath. Together, these ensure that the equilibrium distribution of sampled configurations $\mathbf{r}$ is given by:

$$\pi_{\mathbf{r}}(\mathbf{r}) \propto \exp[-u(\mathbf{r})/(k_B T)]. \quad (71)$$

In practice Equation 70 is discretized with a finite timestep Δt .

Simulation: To obtain training data, we simulate the mini-protein chignolin using Langevin molecular dynamics in OpenMM with the AMBER14 all-atom force field and the TIP3P water model. The chignolin peptide (PDB ID: 1UAO) is solvated in a cubic water box with 1nm padding and neutralized. Simulations are performed at 360K using Langevin dynamics (1ps⁻¹ friction, 2fs timestep) with PME electrostatics and a 1nm cutoff. After energy minimization and velocity initialization, we simulate the system for 1μs and save coordinates every 2ps.

Data preparation: After simulation, we apply preprocessing steps to the raw data. These include removing the water solvent and hydrogen atoms, retaining only the heavy atoms of the protein. We then center the coordinates of each individual configuration and superpose the configurations along the simulated trajectory.

B.1.2 Model Parameterization

Network architecture: To parameterize the flow, we use the $E(3)$ -equivariant graph neural network (GNN) architecture proposed by [Satorras et al., 2021]. Equivariance is achieved using equivariant graph convolutional layers (EGCL). Given coordinates $\mathbf{x}_i^{(l)} \in \mathbb{R}^3$ and node embeddings $\mathbf{h}_i^{(l)} \in \mathbb{R}^{d_H}$ for each node i , the output node features and coordinates of the l -th EGCL layer are computed as follows:

$$\mathbf{m}_{ij} = \varphi_e \left(\mathbf{h}_i^{(l)}, \mathbf{h}_j^{(l)}, \gamma \left(\|\mathbf{x}_i^{(l)} - \mathbf{x}_j^{(l)}\|^2 \right), a_{ij} \right), \quad (72)$$

$$\mathbf{x}_i^{(l+1)} = \mathbf{x}_i^{(l)} + \frac{1}{M_i - 1} \sum_{j \neq i} \left(\mathbf{x}_i^{(l)} - \mathbf{x}_j^{(l)} \right) \varphi_x(\mathbf{m}_{ij}), \quad (73)$$

$$\mathbf{m}_i = \sum_{j \neq i} \mathbf{m}_{ij}, \quad (74)$$

$$\mathbf{h}_i^{(l+1)} = \varphi_h \left(\mathbf{h}_i^{(l)}, \mathbf{m}_i^{(l)} \right). \quad (75)$$

Here, φ denotes functions parameterized by multi-layer perceptrons (MLPs), and γ represents a d_F -dimensional Fourier feature encoding function of the distance between two node coordinates. Furthermore, a_{ij} denotes information associated with the edge between nodes i and j , and M_i denotes the number of nodes in the one-hop neighborhood of node i . The full network consists of L such layers. As initial hidden node embeddings $\mathbf{h}_i^{(0)}$, we use a concatenation of linear embeddings of the particle’s atom type $a \in [0, 1]^{N_A}$, associated bead type $b \in [0, 1]^{N_B}$, and the interpolation time $t \in [0, 1]$. Moreover, we do not include additional edge information a_{ij} .

Noise distribution: For the projected-out atoms, we define a residue-wise target latent distribution. The latent position $\epsilon_{I,i}$ of the i -th atom in the I -th residue is sampled from a Gaussian distribution centered at the position $\mathbf{R}_I$ of the corresponding C_α atom:

$$\epsilon_{I,i} \sim \mathcal{N}(\mathbf{R}_I, \sigma^2 \mathbf{1}), \quad (76)$$

with variance σ^2 .

We report the hyperparameter choices for the model’s architecture and for the noise distribution in Table 2.

Table 2: Architectural hyperparameters of the model trained for backmapping the C_α representation of chignolin. We report the choices for the $E(3)$ -equivariant GNN parameterization of the velocity field and the noise distribution.

Hyperparameter	Value
Number of layers L	6
Number of Fourier features d_F	6
Hidden dimensionality d_H	65
Latent variance σ^2	0.04

B.1.3 Training Details

Split-flows: We train our model on the conditional flow matching objective presented in Section 3.4, using the coupling described in Section 4 and a linear reference interpolant:

$$I : [0, 1] \times \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}^n, \quad t, \mathbf{x}_0, \mathbf{x}_1 \mapsto I_t(\mathbf{x}_0, \mathbf{x}_1) = (1 - t)\mathbf{x}_0 + t\mathbf{x}_1. \quad (77)$$

All training hyperparameters are listed in Table 3. Training takes approximately 40 hours on an NVIDIA A30 GPU with 30 GB of memory.

Table 3: Training hyperparameters of the model trained for backmapping the C_α representation of chignolin.

Hyperparameter	Value
Optimizer	Adam
Learning rate (LR)	3×10^{-4}
LR scheduler	One-cycle
Weight decay	1×10^{-3}
Batch size	64
Number of opt. steps	24,760

TC-VAE: We retrain the TC-VAE [Shmilovich et al., 2022] using the code and hyperparameters provided by the authors. We find that the additional energy regularization proposed by the authors consistently leads to numerical instabilities, resulting in NaN values in the loss. We therefore train the model without the energy regularization. Training takes approximately 120 hours on an NVIDIA A30 GPU with 30 GB of memory.

CG-back: Since CG-back [Torre and Sugita, 2025] is a transferable model, we utilize the pretrained model (M), provided by the authors.

Flow-back: Flow-back [Berglaga et al., 2025] is a transferable model. We hence use the pretrained model (Pretrained) provided by the authors.

B.1.4 Evaluation Metrics

We evaluate and compare the backmapping capabilities of split-flows and reference methods using several metrics. In this section, we provide additional details on their computation. We will denote the reference configurations as $\mathbf{r}$ and $\mathbf{R} = M(\mathbf{r})$ and the reconstructed configurations as $\hat{\mathbf{r}}$ and $\hat{\mathbf{R}} = M(\hat{\mathbf{r}})$ for the fine- and coarse-grained resolutions, respectively.

Wasserstein-1 distance of the internal energy distribution: The internal potential energy of each configuration is computed using the AMBER14 all-atom force field under vacuum conditions. For each configuration $\mathbf{r}$, we add hydrogen atoms using OpenMM and relax the positions via energy minimization (up to 1000 iterations or until the forces fall below $0.1 \text{ kJ mol}^{-1} \text{ nm}^{-1}$). We then compute the Wasserstein-1 distance between the distributions of internal energies of the reference and reconstructed configurations. In Figure 10, we show histograms of the energy distributions.

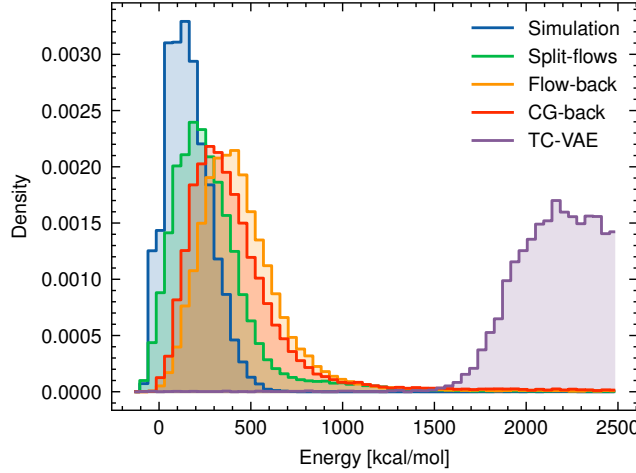


Figure 10: Distributions of internal energies of configurations from the reference simulation and backmapped configurations obtained from the methods in the comparison. The internal energy is evaluated using the AMBER14 all-atom force field.

Coarse-grained RMSD: To quantify the consistency between coarse-grained configurations and their corresponding coarse-grained representatives, we compute the root-mean-squared deviation (RMSD) between the C_α atoms of the coarse-grained and backmapped configurations. The RMSD is well-defined up to global rotations and translations, and can be expressed as:

$$\text{RMSD}_{\text{cg}}(\mathbf{R}, \hat{\mathbf{R}}) = \min_{Q, \mathbf{t}} \left[\frac{1}{N} \sum_{i=1}^N \left\| Q \mathbf{R}_i + \mathbf{t} - \hat{\mathbf{R}}_i \right\|^2 \right]^{1/2}, \quad (78)$$

where $\mathbf{R}_i$ and $\hat{\mathbf{R}}_i$ denote the positions of the i -th coarse-grained particle for the reference and reconstructed configurations, respectively. Furthermore, $Q \in SO(3)$ is a global rotation matrix and $\mathbf{t} \in \mathbb{R}^3$ a global translation vector.

Relative graph-edit distance: We measure topological reconstruction quality using the relative graph edit distance between the molecular graph obtained from inter-atomic distances and the reference graph. Given a reconstructed configuration $\hat{\mathbf{r}}$, we construct an adjacency matrix $\hat{A}$ based on the Van-der-Waals cutoff values c in Table 4, where the entry at position ij is defined as:

$$\hat{A}_{ij} = \begin{cases} 1 & \text{if } \|\hat{\mathbf{r}}_i - \hat{\mathbf{r}}_j\|_2 < s(c_i + c_j), \\ 0 & \text{otherwise} \end{cases}, \quad (79)$$

where c_i and c_j are the respective cutoff values, and $s = 1.3$ is a scaling factor. We then compute the relative

graph edit distance as:

$$D_G = \frac{\sum_{ij} (A - \hat{A})_{ij}}{\sum_{ij} A_{ij}}, \quad (80)$$

where A is the adjacency matrix of the reference molecular graph.

Table 4: VDW cutoff values in nm for atoms with atomic numbers 1–107.

Z	Cutoff	Z	Cutoff	Z	Cutoff	Z	Cutoff	Z	Cutoff	Z	Cutoff	Z	Cutoff	Z	Cutoff
1	0.023	2	0.093	3	0.068	4	0.035	5	0.083	6	0.068	7	0.068	8	0.068
9	0.064	10	0.112	11	0.097	12	0.110	13	0.135	14	0.120	15	0.075	16	0.102
17	0.099	18	0.157	19	0.133	20	0.099	21	0.144	22	0.147	23	0.133	24	0.135
25	0.135	26	0.134	27	0.133	28	0.150	29	0.152	30	0.145	31	0.122	32	0.117
33	0.121	34	0.122	35	0.121	36	0.191	37	0.147	38	0.112	39	0.178	40	0.156
41	0.148	42	0.147	43	0.135	44	0.140	45	0.145	46	0.150	47	0.159	48	0.169
49	0.163	50	0.146	51	0.146	52	0.147	53	0.140	54	0.198	55	0.167	56	0.134
57	0.187	58	0.183	59	0.182	60	0.181	61	0.180	62	0.180	63	0.199	64	0.179
65	0.176	66	0.175	67	0.174	68	0.173	69	0.172	70	0.194	71	0.172	72	0.157
73	0.143	74	0.137	75	0.135	76	0.137	77	0.132	78	0.150	79	0.150	80	0.170
81	0.155	82	0.154	83	0.154	84	0.168	85	0.170	86	0.240	87	0.200	88	0.190
89	0.188	90	0.179	91	0.161	92	0.158	93	0.155	94	0.153	95	0.151	96	0.150
97	0.150	98	0.150	99	0.150	100	0.150	101	0.150	102	0.150	103	0.150	104	0.157
105	0.149	106	0.143	107	0.141										

Fiber diversity: To measure the diversity of generated structures for a given coarse-grained representative, we draw M samples on the fiber $\mathcal{F}(\mathbf{R})$ for a given coarse-grained representative $\mathbf{R} = M(\mathbf{r})$. Following [Jones et al., 2023], we then define a diversity score η_{div} as the ratio of the average pairwise RMSD (see Equation 78) between all generated configurations and the average RMSD between each generated structure and the reference configuration $\mathbf{r}$:

$$\eta_{\text{div}} = \frac{\frac{2}{M(M-1)} \sum_{m \neq k} \text{RMSD}(\mathbf{r}_m, \mathbf{r}_k)}{\frac{1}{M} \sum_m \text{RMSD}(\mathbf{r}_m, \mathbf{r})}, \quad (81)$$

where here $\mathbf{r}_m$ and $\mathbf{r}_k$ denote the m -th and k -th sample on the fiber.

B.2 Solute in a Lipid Bilayer

B.2.1 Data Generation

Simulation: The simulated data is due to [Nagel and Bereau, 2025], who simulate a coarse-grained POPC lipid bilayer interacting with a two-bead C1P3 solute using the Martini 3 force field. Simulations are performed in GROMACS 2024.3 with a time step of 0.02ps and a simulation box of $6 \times 6 \times 10 \text{ nm}^3$ under periodic boundary conditions. The systems are first equilibrated for 200ps in an NPT ensemble at 298K and 1bar. Subsequently, the system is simulated for 1 μ s in an NVT ensemble with a constant biasing force of $10 \text{ kcal mol}^{-1} \text{ nm}^{-1}$ dragging the solute through the simulation box. Frames are saved every 0.2ps.

Data preparation: The simulated configurations are translated such that the membrane center is at the origin. We then extract the positions of the C1 and P3 beads from the simulated trajectory to compute a two-dimensional description consisting of the distance z between the center of mass of the solute and the membrane center, and the relative orientation ϑ with respect to the z -axis.

B.2.2 Model Parameterization

Network architecture: We parameterize the velocity field $\mathbf{v}_t^\theta$ of the split-flow using a simple MLP. To account for the periodicity in the distance from the membrane center $z \in [-\frac{L}{2}, \frac{L}{2}]$, we apply a sine-cosine input parameterization to the MLP:

$$z \mapsto \left[\sin\left(\frac{2\pi}{L}z\right) \quad \cos\left(\frac{2\pi}{L}z\right) \right]^T. \quad (82)$$

Noise distribution: As the target noise distribution $\pi_{\epsilon|R}$, we use a uniform distribution $\mathcal{U}([0, \pi])$ over the angular domain. With this choice, the flow directly provides access to the excess quantities, i.e., the excess local mapping

entropy and the excess information loss.

We give our architectural hyperparameter choices in Table 5.

Table 5: Architectural hyperparameters of the model trained for backmapping the reduced representation of a solute dragged through a lipid bilayer. We report the choices for the MLP parameterization of the velocity field.

Hyperparameter	Value
Number of layers L	3
Hidden dimensionality d_H	32
Activation function	ReLU

B.2.3 Training Details

Training takes approximately 40 minutes on an NVIDIA GeForce RTX 4060 with 8 GB of memory. We report all hyperparameter choices for training the model in Table 6

Table 6: Training hyperparameters of the model trained for backmapping the reduced representation of a solute dragged through a lipid bilayer.

Hyperparameter	Value
Optimizer	Adam
Learning rate (LR)	1×10^{-3}
LR scheduler	–
Weight decay	0
Batch size	2048
Number of opt. steps	195,300

B.2.4 KDE Comparison Details

As a baseline, we fit a kernel density estimator (KDE) to the training data samples $\mathbf{r}$ and $\mathbf{R}$ to obtain estimates π_R^{KDE} and π_r^{KDE} for the fine- and coarse-grained densities, respectively. We then bin the test data samples according to $\mathbf{R} = [z]^\top$ to obtain an estimate of the local mapping entropy in Equation 9. In Figure 6 (B), we show the resulting excess information loss landscape across the simulation box. Furthermore, in Figure 11, we present a correlation plot between the split-flow and KDE estimates of the mapping entropy, which shows great agreement of the two methods, with a mean absolute error of 0.027 and a Pearson correlation coefficient of 0.99.

B.3 Alanine Dipeptide

B.3.1 Data Generation

Simulation: To obtain training data, we simulate alanine dipeptide using Langevin molecular dynamics (see Appendix B.1.1) in OpenMM with the AMBER14 all-atom force field and the TIP3P water model. The alanine dipeptide molecule is solvated in a cubic water box with 1nm padding and neutralized. Simulations are performed at 600K using Langevin dynamics (1ps^{-1} friction, 2fs timestep) with PME electrostatics and a 1nm cutoff. After energy minimization and velocity initialization, we simulate the system for 1 μs and save coordinates every 2ps.

Data preparation: To preprocess the raw simulation data, we remove the water solvent, retaining only the atoms of alanine dipeptide. We then center the coordinates of each individual configuration and superpose the configurations along the simulated trajectory.

B.3.2 Rigidity of the Coarse-Grained Representation

The visualization of the local mapping entropy in the (ϕ, ψ) plane of Ramachandran angles relies on the rigidity of the bonds and angles between the five atoms in the coarse-grained representation. In Table 7, we report the

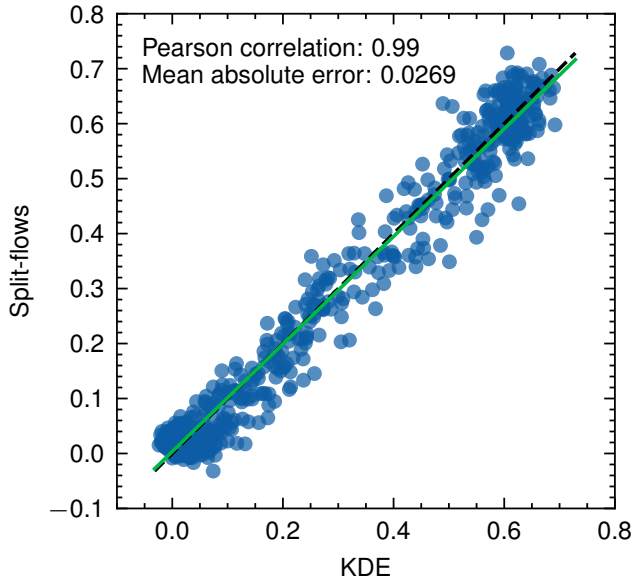


Figure 11: Correlation plot of split-flow and KDE estimates of the local information loss across the lipid membrane. The black dashed line denotes perfect agreement and nearly coincides with a linear fit to the data, shown in green.

mean relative deviations of bond lengths and bond angles. We observe only small bond and angular fluctuations up to approximately 3%, indicating that these degrees of freedom contribute negligibly to the configurational entropy.

Table 7: Relative deviation of the bond lengths and bond angles to the their respective mean values. We report mean and standard deviation evaluated over trajectory used for training.

Bond / Angle	Relative deviation [%]
Bond C–N	2.07 ± 1.56
Bond N–CA	2.20 ± 1.64
Bond CA–C	2.14 ± 1.64
Bond C–N	2.01 ± 1.57
Angle C–N–CA	2.90 ± 2.20
Angle N–CA–C	3.20 ± 2.43
Angle C–C–N	2.73 ± 2.07

B.3.3 Model Parameterization

We parameterize the model analogously to the model trained on chignolin, as presented in Appendix B.1.2. We give our hyperparameter choices in Table 8.

B.3.4 Training Details

Training takes about 18 hours on an NVIDIA A30 GPU with 30 GB of memory. We report all hyperparameter choices for training in Table 9

Table 8: Architectural hyperparameters of the model trained for backmapping the coarse-grained representation of alanine dipeptide. We report the choices for the $E(3)$ -equivariant GNN parameterization of the velocity field and the noise distribution.

Hyperparameter	Value
Number of layers L	4
Number of Fourier features d_F	6
Hidden dimensionality d_H	129
Latent variance σ^2	0.04

Table 9: Training hyperparameters of the model trained for backmapping the coarse-grained representation of alanine dipeptide.

Hyperparameter	Value
Optimizer	Adam
Learning rate (LR)	1×10^{-4}
LR scheduler	Exponential ($\gamma = 0.999$)
Weight decay	0
Batch size	64
Number of opt. steps	50,500