
Bregman contrastive estimation as a general framework for learning unnormalized models: Efficiency, robustness and optimal noise

Yuto Fujii¹

Hiroaki Sasaki²

Takafumi Kanamori^{1,3}

¹Dept. of Math. & Comp. Sci., Institute of Science Tokyo, Tokyo, Japan

²Dept. of Math. Info., Meiji Gakuin University, Kanagawa, Japan

³RIKEN, Center for Advanced Intelligence Project, Tokyo, Japan

Abstract

Learning unnormalized models is an ubiquitous problem in modern statistics and machine learning. Noise contrastive estimation (NCE) takes a popular approach based on solving a binary classification problem: Unnormalized models are learned by discriminating input data with artificially generated noise data. This paper follows this approach, but proposes a general framework based on the Bregman divergence. By selecting the convex function in the divergence, our framework includes existing methods as special cases, and novel variants of NCE can be also derived. Learning unnormalized models in the proposed framework is performed by estimating the posterior probability in the binary classification, i.e., the composition of the logistic function and the ratio of data and noise distributions. Due to the boundedness of the logistic function, our framework is advantageous in outlier-robustness. In fact, we theoretically prove that robust estimation is possible in our framework under a variety of convex functions in the Bregman divergence. Furthermore, the asymptotic efficiency is also investigated, implying a trade-off between efficiency and robustness in our framework. Inspired by these results, we derive the optimal noise distribution under some constraints for robustness. Finally, we numerically demonstrate the robustness of the novel variants of NCE and the derived optimal noise distribution.

1 INTRODUCTION

It is the central problem to learn unnormalized parametric models in modern statistics and machine learning. In contrast with the baseline methods (e.g., maximum likelihood estimation (MLE)), learning unnormalized models does not

require them to be properly normalized, and thus offers a wider range of practical applications. A well-known application is generative models, e.g., generative adversarial networks (GAN) [Goodfellow et al., 2014] and diffusion models [Song and Ermon, 2019, Song et al., 2021]. Another example is self-supervised learning [Liu et al., 2020], which can be regarded as obtaining a useful representation of data through estimation of unnormalized models. Other applications include large language models [Mnih and Teh, 2012, Chen et al., 2024].

A number of practical methods for learning unnormalized models have been proposed so far [Besag, 1975, Hinton, 2002, Hyvärinen, 2005, Gutmann and Hyvärinen, 2012, Takenouchi and Kanamori, 2017, Uehara et al., 2020]. One of the promising methods is *noise contrastive estimation* (NCE) [Gutmann and Hyvärinen, 2012], which forms the basis of GAN and self-supervised learning. NCE learns unnormalized models by discriminating input data from artificially generated noise data through logistic regression, and thus can be seen as solving a binary classification problem. As in MLE [Van der Vaart, 1998, Wasserman, 2004], NCE has a good theoretical property called asymptotic efficiency [Gutmann and Hyvärinen, 2012]. On the other hand, this efficiency of NCE comes at the cost of robustness because logistic regression used in NCE is well-known to be sensitive to the contamination of outliers [Feng et al., 2014].

The crucial point in NCE is how to make noise data samples or choose the noise distribution because the practical performance strongly depends on them. To cope with this problem, Ceylan and Gutmann [2018] proposed to generate noise data by injecting a small amount of additive noise into input data, and demonstrated that the proposed approach works well for data with a lower-dimensional manifold structure (e.g., ring). GAN takes an iterative approach that updates noise samples (i.e., noise distribution) so that it makes difficult or almost impossible to discriminate them from input data samples Goodfellow et al. [2014]. A similar iterative approach has been proposed based on normalizing flow [Gao et al., 2020] where the noise distribution is

updated towards the data distribution. See a recent review paper for NCE [Gutmann et al., 2022] and references therein. Overall, a common belief behind these approaches would be that the data or near-data distribution is an optimal choice for the noise distribution. However, in spite of this belief, it has recently been shown that the optimal noise distribution could be different from the data distribution [Chehab et al., 2022], meaning that there exists a clear gap between them.

This paper follows the same classification approach as NCE yet proposes a general framework to learn unnormalized models with the Bregman divergence [Bregman, 1967]. Previously, similar frameworks have been proposed based on the Bregman divergence, which directly estimate the ratio between data and noise distributions for learning unnormalized models [Gutmann and Hirayama, 2011, Uehara et al., 2020]. In contrast, the proposed framework estimates the posterior probability in the classification problem, i.e., the composition of the logistic function and ratio of the two probability distributions, and thus *not* directly the ratio. This subtle difference makes a great benefit in terms of robustness against outliers. The posterior probability is always bounded even if the ratio is unbounded due to the logistic function. This boundedness is critical for robust estimation because the ratio can be arbitrarily large when outliers move to tails of the data or noise distribution. Previously, an outlier-robust variant of NCE has been proposed [Sasaki and Takenouchi, 2024] but is also a special case of the proposed framework.

After proposing our framework, we perform detailed asymptotic analysis in robustness and efficiency. Our analysis based on the influence function [Huber and Ronchetti, 2009] guarantees the robustness of the proposed framework on a variety of convex functions in the Bregman divergence, while NCE is a rare exception, i.e., it is strongly hampered by the contamination by outliers. The asymptotic efficiency is achievable in the proposed framework as well with a particular choice of the noise distribution in contrast with NCE. Overall, these theoretical results imply a trade-off between efficiency and robustness in our framework.

Subsequently, we establish the optimal noise distribution in the proposed framework. In stark contrast with Chehab et al. [2022], we propose to impose novel constraints to optimize the noise distribution. Interestingly, the optimal noise distribution derived under one of the constraints is a combination of the noise distribution in Chehab et al. [2022] and data distribution, and thus might give an insight to bridge a gap to the aforementioned common belief behind NCE. Furthermore, inspired by the asymptotic analysis, we propose other constraints for robustness, and derive the optimal noise distributions under them with substantially different proofs from that in Chehab et al. [2022]. Finally, we experimentally show the robustness of the derived optimal noise distribution and novel variants of NCE in the proposed framework.

This paper is organized as follows: Section 2 formulates the

problem in this paper and reviews noise contrastive estimation and existing frameworks for learning unnormalized models. Our framework based on the Bregman divergence is proposed in Section 3. Section 4 performs asymptotic analysis in terms of efficiency and robustness, while Section 5 derives the optimal noise distribution in the proposed framework. The novel variants of NCE in our framework and the optimal noise distribution are experimentally investigated in Section 6. Section 7 concludes this paper. A notation table is given in Appendix A.

2 BACKGROUND

We first suppose that there exists a σ -finite reference measure space $(\Omega, \mathcal{B}, \mu)$, and the probability measure P_d for data is absolutely continuous with respect to μ , i.e., it has the density as $p_d = \frac{dP_d}{d\mu}$. Then, we assume that n samples of data vectors $\mathbf{x}_i$ drawn from $p_d(\mathbf{x})$ are available as

$$\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n \stackrel{\text{iid}}{\sim} p_d(\mathbf{x}).$$

Then, this paper considers the problem of estimating the parameter vector $\boldsymbol{\theta}$ in a parametric model,

$$p_m(\mathbf{x}; \boldsymbol{\theta}) = \frac{q_m(\mathbf{x}; \boldsymbol{\theta})}{Z(\boldsymbol{\theta})},$$

from $\mathcal{X}$ where $Z(\boldsymbol{\theta})$ is the partition function, and $q_m(\mathbf{x}; \boldsymbol{\theta})$ is some non-negative function with $\boldsymbol{\theta}$. This section reviews noise contrastive estimation and an existing framework for learning unnormalized models based on the Bregman divergence.

2.1 NOISE CONTRASTIVE ESTIMATION

Noise contrastive estimation (NCE) learns a parametric model by solving a binary classification problem where the following dataset is supposed to be additionally available:

$$\mathcal{Y} = \{\mathbf{y}_i\}_{i=1}^{n_y} \stackrel{\text{iid}}{\sim} p_n(\mathbf{y}),$$

where $\mathbf{y}_i$ and p_n are called the *noise samples* and *noise distribution*, respectively¹. p_n is fundamentally required to have two properties: (i) the functional form of p_n is known and (ii) sampling from p_n is possible. By regarding $\mathcal{X}$ and $\mathcal{Y}$ as datasets from binary classes, the conditional distribution of data $\mathbf{u}$ given class label C is defined as

$$p(\mathbf{u}|C=0) = p_d(\mathbf{u}) \quad \text{and} \quad p(\mathbf{u}|C=1) = p_n(\mathbf{u}).$$

¹As in p_d , we assume that the probability measure P_n for noise data is absolutely continuous with respect to μ and has the density as $p_n = \frac{dP_n}{d\mu}$.

Then, the posterior probability is formulated as

$$\begin{aligned} p(C=0|\mathbf{u}) &= \frac{p_d(\mathbf{u})}{p_d(\mathbf{u}) + \nu p_n(\mathbf{u})} = \frac{r(\mathbf{u})}{r(\mathbf{u}) + \nu} \\ p(C=1|\mathbf{u}) &= \frac{\nu p_n(\mathbf{u})}{p_d(\mathbf{u}) + \nu p_n(\mathbf{u})} = \frac{\nu}{r(\mathbf{u}) + \nu}, \end{aligned} \quad (1)$$

where $\nu := \frac{p(C=1)}{p(C=0)}$ with the class probability $p(C)$ and $r(\mathbf{u}) := p_d(\mathbf{u})/p_n(\mathbf{u})$ is the ratio of two probability distributions.

Learning in NCE is performed through the posterior estimation. To this end, NCE expresses the partition function as a single parameter c and defines an unnormalized model $p_m^0(\mathbf{u}; \alpha)$ as

$$\log p_m^0(\mathbf{u}; \alpha) = \log q_m(\mathbf{u}; \theta) - c.$$

Then, all parameters $\alpha = (\theta, c)$ are estimated by fitting the following model to the true posterior probability:

$$\begin{aligned} f_0(\mathbf{u}; \alpha) &= \frac{p_m^0(\mathbf{u}; \alpha)}{p_m^0(\mathbf{u}; \alpha) + \nu p_n(\mathbf{u})} = \frac{r_m(\mathbf{u}; \alpha)}{r_m(\mathbf{u}; \alpha) + \nu} \\ f_1(\mathbf{u}; \alpha) &= \frac{\nu p_n(\mathbf{u})}{p_m^0(\mathbf{u}; \alpha) + \nu p_n(\mathbf{u})} = \frac{\nu}{r_m(\mathbf{u}; \alpha) + \nu}, \end{aligned} \quad (2)$$

where $r_m(\mathbf{x}; \alpha) := p_m^0(\mathbf{u}; \alpha)/p_n(\mathbf{u})$. Then, the risk function of NCE is formulated through logistic regression as

$$\mathcal{L}_{\text{lr}}(\alpha) = -\mathbb{E}_d[\log f_0(\mathbf{X}; \alpha)] - \nu \mathbb{E}_n[\log f_1(\mathbf{Y}; \alpha)],$$

where $\mathbb{E}_d$ and $\mathbb{E}_n$ denote the expectations over p_d and p_n , respectively. In practice, the empirical version of $\mathcal{L}_{\text{lr}}(\alpha)$ is minimized to estimate the parameter vector α .

2.2 DIRECT RATIO ESTIMATION WITH BREGMAN DIVERGENCE

Learning unnormalized models is possible by estimating the ratio $r_m(\mathbf{x}; \alpha) = p_m^0(\mathbf{x}; \alpha)/p_n(\mathbf{x})$ if p_n is known as assumed in NCE. In fact, (2) indicates that NCE can be regarded as estimating $r_m(\mathbf{x}; \alpha)$ through logistic regression. Based on the ratio estimation, a general framework for learning unnormalized models has been proposed [Gutmann and Hirayama, 2011, Uehara et al., 2020] with the (separable) Bregman divergence, whose risk is defined by

$$\begin{aligned} & -\mathbb{E}_d[\varphi'(r_m(\mathbf{X}; \alpha))] \\ & + \mathbb{E}_n[\{\varphi'(r_m(\mathbf{Y}; \alpha)) r_m(\mathbf{Y}; \alpha) - \varphi(r_m(\mathbf{Y}; \alpha))\}], \end{aligned}$$

where $\varphi(\cdot)$ is a strictly convex function on $\mathbb{R}_+$ and $\varphi'(z) := \frac{d}{dz} \varphi(z)$. By selecting φ , a variety of risk functions can be derived.

This line of work is similar to our approach, but there exists a fundamental difference: We do not directly estimate the ratio $r(\mathbf{x})$, but the posterior probability in (1), i.e., the composition of the logistic function and ratio $r(\mathbf{x})$. The notable

point is that the posterior probability is always bounded due to the composition by the logistic function, while ratio $r(\mathbf{x})$ can be unbounded. This subtle point makes robust learning possible when outliers lie on the tails of the data or noise distribution. Indeed, it has been theoretically shown that direct ratio estimation can be strongly hampered by the contamination of outliers even when robust divergences are employed [Kanamori and Sugiyama, 2014, Sasaki and Takenouchi, 2024].

3 BREGMAN CONTRASTIVE ESTIMATION

The main idea in our framework is to employ the Bregman divergence between the posterior probability $p(C|\mathbf{u})$ and its model $f_C(\mathbf{u}; \alpha)$ for which the risk function is defined as

$$\int p(\mathbf{u}) \sum_{C \in \{0,1\}} [\varphi(p(C|\mathbf{u})) - \varphi(f_C(\mathbf{u}; \alpha)) - \varphi'(f_C(\mathbf{u}; \alpha))(p(C|\mathbf{u}) - f_C(\mathbf{u}; \alpha))] d\mathbf{u}. \quad (3)$$

Then, by expressing the marginal distribution as

$$p(\mathbf{u}) = \sum_{C \in \{0,1\}} p(\mathbf{u}|C)p(C) \propto p_d(\mathbf{u}) + \nu p_n(\mathbf{u}),$$

Appendix B shows that our risk function can be written by

$$\begin{aligned} \mathcal{L}_{\varphi}^{\text{SB}}(\alpha) &:= -\mathbb{E}_d[\Phi_{\varphi}(\mathbf{X}; \alpha) + \Psi_{\varphi}(\mathbf{X}; \alpha)f_1(\mathbf{X}; \alpha)] \\ &\quad - \nu \mathbb{E}_n[\Phi_{\varphi}(\mathbf{Y}; \alpha) - \Psi_{\varphi}(\mathbf{Y}; \alpha)f_0(\mathbf{Y}; \alpha)], \end{aligned} \quad (4)$$

where $\Phi_{\varphi}(\mathbf{u}; \alpha) := \varphi(f_0(\mathbf{u}; \alpha)) + \varphi(f_1(\mathbf{u}; \alpha))$ and $\Psi_{\varphi}(\mathbf{u}; \alpha) := \varphi'(f_0(\mathbf{u}; \alpha)) - \varphi'(f_1(\mathbf{u}; \alpha))$. In practice, $\mathcal{L}_{\varphi}^{\text{SB}}(\alpha)$ is empirically approximated from data samples as

$$\begin{aligned} \hat{\mathcal{L}}_{\varphi}^{\text{SB}}(\alpha) &:= -\frac{1}{n} \sum_{i=1}^n [\Phi_{\varphi}(\mathbf{x}_i; \alpha) + \Psi_{\varphi}(\mathbf{x}_i; \alpha)f_1(\mathbf{x}_i; \alpha)] \\ &\quad - \frac{1}{n} \sum_{i=1}^{n_y} [\Phi_{\varphi}(\mathbf{y}_i; \alpha) - \Psi_{\varphi}(\mathbf{y}_i; \alpha)f_0(\mathbf{y}_i; \alpha)], \end{aligned}$$

where ν is substituted by $\hat{\nu} := n_y/n$. The parameter vector α is estimated by minimizing $\hat{\mathcal{L}}_{\varphi}^{\text{SB}}(\alpha)$ as

$$\hat{\alpha} = \underset{\alpha}{\operatorname{argmin}} \hat{\mathcal{L}}_{\varphi}^{\text{SB}}(\alpha) = (\hat{\theta}, \hat{c}).$$

We call this method as the *Bregman contrastive estimation* (BCE).

Next, we show that BCE includes NCE and its robust variant called the power NCE (PNCE) [Sasaki and Takenouchi, 2024] as special cases:

Example 1 (NCE). When $\varphi(z) = z \log z$, $\mathcal{L}_{\varphi}^{\text{SB}}(\alpha)$ is reduced to $\mathcal{L}_{\text{lr}}(\alpha)$.

Table 1: Examples of BCE under the choices of φ .

	$\varphi(z)$	risk	reference
NCE	$z \log z$	$\mathcal{L}_{\text{lr}}(\alpha)$	[Gutmann and Hyvärinen, 2012]
PNCE	$\frac{z^{\beta+1}}{\beta(\beta+1)}$	$\mathcal{L}_{\beta}(\alpha)$	[Sasaki and Takenouchi, 2024]
ρ -NCE	$-\log(z + \rho)$	$\mathcal{L}_{\rho}(\alpha)$	This work
τ -NCE	$(z + \tau) \log \frac{z+\tau}{1+\tau}$	$\mathcal{L}_{\tau}(\alpha)$	This work

Example 2 (PNCE). When $\varphi(z) = z^{\beta+1}/(\beta(\beta+1))$, $\mathcal{L}_{\varphi}^{\text{SB}}(\alpha)$ is reduced to the following risk function proposed in Sasaki and Takenouchi [2024]:

$$\begin{aligned} \mathcal{L}_{\beta}(\alpha) &= -\frac{1}{\beta} (\mathbb{E}_{\text{d}}[f_0(\mathbf{X}; \alpha)^{\beta}] + \nu \mathbb{E}_{\text{n}}[f_1(\mathbf{Y}; \alpha)^{\beta}]) \\ &+ \frac{1}{\beta+1} \sum_{c=0}^1 (\mathbb{E}_{\text{d}}[f_c(\mathbf{X}; \alpha)^{\beta+1}] + \nu \mathbb{E}_{\text{n}}[f_c(\mathbf{Y}; \alpha)^{\beta+1}]). \end{aligned}$$

Next, two novel variants of NCE are derived from BCE. The first choice is $\varphi(z) = -\log(z + \rho)$, which has been used in boosting [Murata et al., 2004] and introduced to cope with a mislabeling problem [Takenouchi and Eguchi, 2004]. The second choice can be seen as a generalization of the first one.

Example 3 (ρ -NCE). When $\varphi(z) = -\log(z + \rho)$ with a parameter $\rho > 0$, $\mathcal{L}_{\varphi}^{\text{SB}}(\alpha)$ is reduced to

$$\begin{aligned} \mathcal{L}_{\rho}(\alpha) &= \mathbb{E}_{\text{d}} \left[-\log(f_0(\mathbf{X}; \alpha)f_1(\mathbf{X}; \alpha) + \rho^2 + \rho) \right. \\ &\quad \left. - \frac{f_1(\mathbf{X}; \alpha)\{f_0(\mathbf{X}; \alpha) - f_1(\mathbf{X}; \alpha)\}}{f_0(\mathbf{X}; \alpha)f_1(\mathbf{X}; \alpha) + \rho^2 + \rho} \right] \\ &- \nu \mathbb{E}_{\text{n}} \left[-\log(f_0(\mathbf{Y}; \alpha)f_1(\mathbf{Y}; \alpha) + \rho^2 + \rho) \right. \\ &\quad \left. - \frac{f_0(\mathbf{Y}; \alpha)\{f_0(\mathbf{Y}; \alpha) - f_1(\mathbf{Y}; \alpha)\}}{f_0(\mathbf{Y}; \alpha)f_1(\mathbf{Y}; \alpha) + \rho^2 + \rho} \right]. \end{aligned}$$

We call the learning method based on $\mathcal{L}_{\rho}(\alpha)$ as ρ -NCE.

Example 4 (τ -NCE). When $\varphi(z) = (z + \tau) \log \frac{z+\tau}{1+\tau}$ for $\tau > 0$, $\mathcal{L}_{\varphi}^{\text{SB}}(\alpha)$ is reduced to

$$\begin{aligned} \mathcal{L}_{\tau}(\alpha) &= -\mathbb{E}_{\text{d}} [(1 + \tau) \log f_0^{\tau}(\mathbf{X}; \alpha) + \tau \log f_1^{\tau}(\mathbf{X}; \alpha)] \\ &- \nu \mathbb{E}_{\text{n}} [(1 + \tau) \log f_1^{\tau}(\mathbf{Y}; \alpha) + \tau \log f_0^{\tau}(\mathbf{Y}; \alpha)], \end{aligned}$$

where $f_C^{\tau}(\mathbf{x}; \alpha) := \frac{f_C(\mathbf{x}; \alpha) + \tau}{1 + \tau}$ for $C = 0, 1$. We call the learning method based on $\mathcal{L}_{\tau}(\alpha)$ as τ -NCE.

It is shown theoretically and empirically that τ -NCE and ρ -NCE are robust against outliers. Table 1 summarizes how each method corresponds to a particular choice of φ .

3.1 MODEL SELECTION IN BCE

The selection of the hyper-parameters (e.g., ρ and τ) is important in practice. Here, by following a simple heuristic

approach proposed in Minami and Eguchi [2003] and Sasaki and Takenouchi [2024], a method for model selection based on cross-validation is presented as follows: Taking τ -NCE as an example,

1. Choose a fixed parameter τ_0 called the *anchor parameter* and candidates of τ denoted by T .
2. Divide $\mathcal{X} = \{\mathbf{x}_i\}_{i=1}^n$ into K disjoint subsets $\{\mathcal{X}_k\}_{k=1}^K$.
3. Obtain the estimator $\hat{\alpha}_{\tau}^{(k)}$ for $\tau \in T$ from $\mathcal{X}$ without $\mathcal{X}_k$ and reference data vectors $\mathbf{y}_j$ which are generated by keeping the same sample ratio $\hat{\nu}$.
4. Compute $\hat{\mathcal{L}}_{\tau_0}$ with the anchor parameter τ_0 from $\mathcal{X}_k$ as $\text{CV}(k, \tau) := \hat{\mathcal{L}}_{\tau_0}(\hat{\alpha}_{\tau}^{(k)})$.
5. Choose the optimal value of τ on T that minimizes $\frac{1}{K} \sum_{k=1}^K \text{CV}(k, \tau)$.

It seems to still matter the choice of the anchor parameter in cross-validation, but our empirical results show that this method reasonably works well on a range of the anchor parameters.

4 ASYMPTOTIC ANALYSIS: ROBUSTNESS AND EFFICIENCY

This section performs asymptotic analysis to BCE in terms of robustness and efficiency. Throughout this section, let us assume that for all θ ,

$$\mathcal{D} := \text{supp}(p_{\text{m}}(\cdot; \theta)) \subseteq \text{supp}(p_{\text{n}}), \quad (5)$$

where $\text{supp}(p_{\text{m}}(\cdot; \theta)) := \{\mathbf{x} \mid p_{\text{m}}(\mathbf{x}; \theta) \neq 0\}$ and $\text{supp}(p_{\text{n}}) := \{\mathbf{y} \mid p_{\text{n}}(\mathbf{y}) \neq 0\}$. This support assumption is standard in the analysis of NCE [Gutmann and Hyvärinen, 2012, Matsuda et al., 2021, Sasaki and Takenouchi, 2024], and required to estimate $p_{\text{m}}(\cdot; \theta)$ over the full support. Furthermore, we assume that $p_{\text{m}}^0(\mathbf{x}; \alpha)$ is differentiable with respect to α such that the gradient $\mathbf{g}_{\text{m}}(\mathbf{x}; \alpha) := \nabla_{\alpha} \log p_{\text{m}}^0(\mathbf{x}; \alpha)$ is properly computed.

4.1 ROBUSTNESS OF BCE

Here, the robustness of BCE is investigated based on the *influence function* [Huber and Ronchetti, 2009, Hampel et al., 2011, Maronna et al., 2019]. To this end, we start by

modeling the contaminated probability distributions as

$$\bar{p}_d(\mathbf{x}) = (1 - \varepsilon)p_d(\mathbf{x}) + \varepsilon\delta_{\mathbf{x}_o}(\mathbf{x}), \quad (6)$$

where $0 \leq \varepsilon < 1$ is a constant, $\mathbf{x}_o$ is an outlier, and $\delta_{\mathbf{x}_o}(\cdot)$ denotes the Dirac measure centered at $\mathbf{x}_o$ and is sufficient for the worse-case local analysis. As the contamination ratio ε is increased, the dataset is contaminated by more outliers. Based on the contaminated probability distribution (6), we define the following estimators:

$$\alpha_\star = \operatorname{argmin}_{\alpha} \mathcal{L}_{\varphi}^{\text{SB}}(\alpha) \text{ and } \alpha_\varepsilon = \operatorname{argmin}_{\alpha} \bar{\mathcal{L}}_{\varphi}^{\text{SB}}(\alpha), \quad (7)$$

where by expressing the expectation over $\bar{p}_d(\cdot)$ as $\bar{\mathbb{E}}_d$,

$$\begin{aligned} \bar{\mathcal{L}}_{\varphi}^{\text{SB}}(\alpha) := & -\bar{\mathbb{E}}_d [\Phi(\mathbf{X}; \alpha) + \Psi(\mathbf{X}; \alpha)f_1(\mathbf{X}; \alpha)] \\ & - \nu \mathbb{E}_n [\Phi(\mathbf{Y}; \alpha) - \Psi(\mathbf{Y}; \alpha)f_0(\mathbf{Y}; \alpha)]. \end{aligned}$$

Based on these estimators, the influence function (IF) is defined in the limit of $\varepsilon \rightarrow 0$ as follows:

$$\text{IF}(\mathbf{x}_o) = \lim_{\varepsilon \rightarrow 0} \frac{\alpha_\varepsilon - \alpha_\star}{\varepsilon}. \quad (8)$$

IF($\mathbf{x}_o$) helps us to understand how the asymptotic estimator $\alpha_\star$ is influenced by infinitesimal contamination of outliers. A useful notion based on IF is *B-robustness*: If $\sup_{\mathbf{x} \in \mathcal{D}} \|\text{IF}(\mathbf{x})\| < \infty$, $\alpha_\star$ is said to be B-robust. B-robustness ensures that the influence from outliers $\mathbf{x}_o$ is limited even if they go far away from the center of data.

We first derive IF and show a condition for B-robustness in the following proposition, and then discuss how mild the condition is:

Proposition 1. *Assume that the Hessian matrix $\mathbf{H}_{\varphi} = \nabla_{\alpha}^2 \mathcal{L}_{\varphi}^{\text{SB}}(\alpha)|_{\alpha=\alpha_\star}$ is invertible. Then, the influence function is given by*

$$\begin{aligned} \text{IF}(\mathbf{x}_o) = & \mathbf{H}_{\varphi}^{-1} \{ \mathbb{E}_d [w_{\varphi}(\mathbf{X}; \alpha_\star)f_1(\mathbf{X}; \alpha_\star)\mathbf{g}_m(\mathbf{X}; \alpha_\star)] \\ & - w_{\varphi}(\mathbf{x}_o; \alpha_\star)f_1(\mathbf{x}_o; \alpha_\star)\mathbf{g}_m(\mathbf{x}_o; \alpha_\star) \}, \quad (9) \end{aligned}$$

where with $\varphi''(z) := \frac{d^2}{dz^2} \varphi(z)$,

$$\begin{aligned} w_{\varphi}(\mathbf{x}; \alpha) := & \{ \varphi''(f_0(\mathbf{x}; \alpha)) + \varphi''(f_1(\mathbf{x}; \alpha)) \} f_0(\mathbf{x}; \alpha)f_1(\mathbf{x}; \alpha). \quad (10) \end{aligned}$$

Furthermore, $\alpha_\star$ is B-robust if the following inequality holds:

$$\sup_{\mathbf{x} \in \mathcal{D}} w_{\varphi}(\mathbf{x}; \alpha_\star) \|\mathbf{g}_m(\mathbf{x}; \alpha_\star)\| < \infty. \quad (11)$$

The proof is deferred to Appendix C. The first term on the right-hand side of (9) is a constant bias and cannot be eliminated but has no influence from outliers. The second term is influenced by outlier $\mathbf{x}_o$, and thus $\alpha_\star$ is guaranteed

to be B-robust if this term is finite, i.e., condition (11) is fulfilled.

For NCE, it is very unlikely that condition (11) is satisfied because the choice of $\varphi(z) = z \log z$ in Example 1 yields a constant weight $w_{\varphi}(\mathbf{x}; \alpha_\star) = 1$. Thus, the left-hand side on (11) becomes $\sup_{\mathbf{x} \in \mathcal{D}} \|\mathbf{g}_m(\mathbf{x}; \alpha_\star)\|$, which can be unbounded in general. Non-robustness of NCE has been previously reported in Sasaki and Takenouchi [2024] as well. Appendix D illustrates that weights $w_{\varphi}(\mathbf{x}; \alpha_\star)$ for PNCE, τ -NCE and ρ -NCE are non-constant and approach to zero if $f_0 \rightarrow 0$ or $f_1 \rightarrow 1$.

Next, we show when condition (11) is fulfilled. To enhance the interpretability, we re-express condition (11) in terms of the ratio model $r_m(\mathbf{x}; \alpha) = p_m^0(\mathbf{x}; \alpha)/p_n(\mathbf{x}; \alpha)$ in the following corollary where $\mathcal{C}^k$ denotes the space of continuous and k -times differentiable functions:

Corollary 1. *Assume that $w_{\varphi}(\mathbf{x}; \alpha_\star)$ is not a constant function and φ belongs to $\mathcal{C}^2$. Then, (11) is equivalent to*

$$\sup_{\mathbf{x} \in \mathcal{D}} \frac{\nu r_m(\mathbf{x}; \alpha_\star)}{(\nu + r_m(\mathbf{x}; \alpha_\star))^2} \|\mathbf{g}_m(\mathbf{x}; \alpha_\star)\| < \infty. \quad (12)$$

The proof is a direct consequence from the finiteness of φ'' because φ is a $\mathcal{C}^2$ function and the domain of φ is $[0, 1]$. It is useful to consider a common situation (e.g., Gaussian distributions) where $\|\mathbf{g}_m(\mathbf{x}; \alpha_\star)\|$ diverges when $p_m^0(\mathbf{x}; \alpha_\star) \rightarrow 0$ as $\mathbf{x}$ goes to the boundary of $\mathcal{D}$. In this situation, $r_m(\mathbf{x}; \alpha)$ approaches zero if $p_n(\mathbf{x})$ has a wider support and slower decay than $p_m^0(\mathbf{x}; \alpha_\star)$ because the growth speed of $\mathbf{g}_m(\mathbf{x}; \alpha_\star) = \nabla_{\alpha} \log p_m^0(\mathbf{x}; \alpha)|_{\alpha=\alpha_\star}$ is logarithmic. Thus, condition (12) seems very mild and could be satisfied under a variety of strictly convex functions φ because the fundamental condition is that φ is a $\mathcal{C}^2$ function. Moreover, a large value of ν would also make condition (12) more easily satisfied, and thus improve the robustness of BCE.

4.2 EFFICIENCY OF BCE

Efficiency of statistical estimators is a classical yet important notion. An estimator is said to be *asymptotically efficient* if the asymptotic variance is equal to the Cramér-Rao lower bound [Lehmann and Casella, 1998], which is the inverse of the Fisher information matrix² defined as

$$\mathbf{I}_F := \mathbb{E}_d [\mathbf{g}_m(\mathbf{X}; \alpha_\star) \mathbf{g}_m(\mathbf{X}; \alpha_\star)^\top].$$

To investigate the efficiency, we first establish the asymptotic normality of BCE and then show some differences to NCE. The consistency of BCE is proved in Appendix E.

²More formally, the Fisher information matrix $\mathbf{I}_F$ is defined with the expectation over $p_m(\cdot, \theta)$. However, Assumptions (A3-4) imply $p_m(\cdot; \theta_\star) = p_d(\cdot)$. Therefore, we defined $\mathbf{I}_F$ in terms of the expectation over $p_d(\cdot)$.

For asymptotic normality, we make the following regularity assumptions:

- (A1) The parameter space Θ of θ is compact.
- (A2) For all θ , $\text{supp } p_m(\cdot; \theta) \subset \text{supp } p_n$.
- (A3) $p_m(\mathbf{x}; \theta)$ is identifiable, i.e., $p_m(\mathbf{x}; \theta_1) = p_m(\mathbf{x}; \theta_2)$ implies $\theta_1 = \theta_2$.
- (A4) $p_d(\mathbf{x})$ belongs to the same parametric family as $p_m(\mathbf{x}; \theta)$, i.e., there exists a constant vector $\theta_0 \in \text{Int}(\Theta)$ (i.e., interior of Θ) such that $p_d(\mathbf{x}) = p_m(\mathbf{x}; \theta_0)$.
- (A5) For all $\epsilon > 0$, $\inf_{\alpha: \|\alpha - \alpha_*\| \geq \epsilon} \mathcal{L}_\varphi^{\text{SB}}(\alpha) > \mathcal{L}_\varphi^{\text{SB}}(\alpha_*)$.
- (A6) $\sup_{\alpha} |\mathcal{L}_\varphi^{\text{SB}}(\alpha) - \hat{\mathcal{L}}_\varphi^{\text{SB}}(\alpha)| \rightarrow 0$ in probability.

Assumptions (A1-4) are quite standard. Assumption (A5) means that α_* is a well-separated point of minimum of $\mathcal{L}_\varphi^{\text{SB}}(\alpha)$ [Van der Vaart, 1998], and Assumption (A6) is the well-known uniform law of large numbers. By definition of α_* and Assumption (A4), we have $p_d(\mathbf{x}) = p_m(\mathbf{x}; \theta_*)$. Therefore, the identifiability condition (A3) ensures that $\theta_0 = \theta_*$, which enables us to express $f_1(\mathbf{u}; \alpha_*)$ as

$$f_1^*(\mathbf{u}) := f_1(\mathbf{u}; \alpha_*) = \frac{\nu p_n(\mathbf{u})}{p_d(\mathbf{u}) + \nu p_n(\mathbf{u})}. \quad (13)$$

In addition, as in Gutmann and Hyvärinen [2012], Matsuda et al. [2021], we suppose the asymptotics that $\hat{\nu} = n_y/n \rightarrow \nu$ as $n \rightarrow \infty$ and $n_y \rightarrow \infty$.

We first establish the asymptotic normality of BCE in the following theorem proved in Appendix F:

Theorem 1 (Asymptotic normality). *Suppose that Assumptions (A1-6) hold. We further assume that there exists an integrable function $b(\mathbf{x})$ such that*

$$|\partial_{ij} p_m^0(\mathbf{x}; \alpha)| \leq b(\mathbf{x}) \text{ and } |\partial_{ijk} p_m^0(\mathbf{x}; \alpha)| \leq b(\mathbf{x}),$$

where $\partial_{ij} := \frac{\partial^2}{\partial \alpha_i \partial \alpha_j}$ and $\partial_{ijk} := \frac{\partial^3}{\partial \alpha_i \partial \alpha_j \partial \alpha_k}$. Then, $\sqrt{n}(\hat{\alpha} - \alpha_*)$ converges to the following normal distribution:

$$\sqrt{n}(\hat{\alpha} - \alpha_*) \xrightarrow{D} N(\mathbf{0}, \mathbf{\Lambda}_\varphi),$$

where $\mathbf{\Lambda}_\varphi := \mathbf{J}_\varphi^{-1} \mathbf{I}_\varphi \mathbf{J}_\varphi^{-1}$, and with $\mathbf{g}_m^w(\mathbf{x}; \alpha) := w_\varphi(\mathbf{x}; \alpha) \mathbf{g}_m(\mathbf{x}; \alpha)$,

$$\begin{aligned} \mathbf{I}_\varphi &= \mathbb{E}_d [f_1^*(\mathbf{X}) \mathbf{g}_m^w(\mathbf{X}; \alpha_*) \mathbf{g}_m^w(\mathbf{X}; \alpha_*)^\top] - \left(1 + \frac{1}{\nu}\right) \\ &\quad \times \mathbb{E}_d [f_1^*(\mathbf{X}) \mathbf{g}_m^w(\mathbf{X}; \alpha_*)] \mathbb{E}_d [f_1^*(\mathbf{X}) \mathbf{g}_m^w(\mathbf{X}; \alpha_*)]^\top \end{aligned} \quad (14)$$

$$\mathbf{J}_\varphi = \mathbb{E}_d [f_1^*(\mathbf{X}) w_\varphi(\mathbf{X}; \alpha_*) \mathbf{g}_m(\mathbf{X}; \alpha_*) \mathbf{g}_m(\mathbf{X}; \alpha_*)^\top]. \quad (15)$$

The asymptotic variance $\mathbf{\Lambda}_\varphi$ of BCE takes a complex form and clearly depends on the strictly convex function φ . The following theorem asserts that NCE is in fact the optimal choice of φ in the sense that it is the closest to the Cramér-Rao bound:

Theorem 2 (Optimality of NCE). *For all strictly convex functions φ , it holds*

$$\mathbf{\Lambda}_\varphi \succeq \mathbf{\Lambda}_{\text{NCE}}, \quad (16)$$

where $\mathbf{\Lambda}_\varphi \succeq \mathbf{\Lambda}_{\text{NCE}}$ means that $\mathbf{\Lambda}_\varphi - \mathbf{\Lambda}_{\text{NCE}}$ is positive semidefinite, and $\mathbf{\Lambda}_{\text{NCE}}$ denotes the asymptotic variance of NCE and is given by (33). Equality in (16) holds when w_φ is a constant function.

The proof is deferred in Appendix G. Theorem 2 indicates the optimality of NCE with respect to the choice of φ . On the other hand, as shown in Section 4.1, NCE is not robust against outliers. This may imply a trade-off in terms of the choice of φ : An efficient choice of φ comes at a cost of losing outlier-robustness.

NCE has been shown to achieve the Cramér-Rao bound in the limit of $\nu \rightarrow \infty$ regardless of the choice of the noise distribution p_n [Gutmann and Hyvärinen, 2012]. In contrast, the following theorem indicates that the Cramér-Rao bound is attained in BCE under a particular choice of p_n in the asymptotic regime as for ν :

Proposition 2 (Noise limit). *Assume that φ belongs to $\mathcal{C}^2$. Then, as $\nu \rightarrow \infty$,*

$$\mathbf{\Lambda}_\varphi = \tilde{\mathbf{J}}^{-1} \tilde{\mathbf{I}} \tilde{\mathbf{J}}^{-1} + O(\nu^{-2}),$$

where

$$\begin{aligned} \tilde{\mathbf{I}} &= \mathbb{E}_d \left[\left(\frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} \right)^2 \mathbf{g}_m(\mathbf{X}; \alpha_*) \mathbf{g}_m(\mathbf{X}; \alpha_*)^\top \right] \\ &\quad - \mathbb{E}_d \left[\frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} \mathbf{g}_m(\mathbf{X}; \alpha_*) \right] \mathbb{E}_d \left[\frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} \mathbf{g}_m(\mathbf{X}; \alpha_*) \right]^\top \\ \tilde{\mathbf{J}} &= \mathbb{E}_d \left[\frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} \mathbf{g}_m(\mathbf{X}; \alpha_*) \mathbf{g}_m(\mathbf{X}; \alpha_*)^\top \right]. \end{aligned}$$

The proof is given in Appendix H. In the limit of $\nu \rightarrow \infty$, the striking features of $\mathbf{\Lambda}_\varphi$ are twofold: (i) $\mathbf{\Lambda}_\varphi$ does not depend on the choice of φ , and (ii) unlike NCE, the noise distribution p_n still exists in $\mathbf{\Lambda}_\varphi$. This possibly indicates that parameter estimation by BCE is more dependent on the choice of p_n than φ in a large noise sample regime.

Proposition 2 shows that $\mathbf{\Lambda}_\varphi$ asymptotically equals to the inverse of $\mathbf{I}_F$, i.e., BCE achieves the Cramér-Rao bound when $p_n = p_d$. However, this choice of p_n makes the weight function w_φ in (10) to be constant, and thus condition (11) for B-robustness is never satisfied when $\|\mathbf{g}_m(\mathbf{x}; \alpha_*)\|$ diverges. This again may reveal a trade-off between efficiency and robustness, but this trade-off seems to be slightly different from the aforementioned one after Theorem 2: The choice of p_n is more relevant in the asymptotic regime (i.e., $\nu \rightarrow \infty$), while the choice of φ would matter in non-asymptotic regime as implied in Theorem 2.

5 THE OPTIMAL NOISE DISTRIBUTION WITH ROBUSTNESS

A common criterion to measure the statistical performance is *mean-squared error* (MSE) between $\hat{\theta}$ and θ_* . As in Chehab et al. [2022], we derive the optimal noise distribution based on the following approximation of MSE:

$$\text{MSE} := \frac{\text{Tr}(\mathbf{\Lambda}_{\varphi})}{n}. \quad (17)$$

This approximation would be valid when n is large because the bias term in MSE often more quickly goes to zero as n increases. In this section, we express $\mathbf{g}_m(\mathbf{x}) := \nabla_{\theta} \log p_m(\mathbf{x}; \theta)|_{\theta=\theta_*}$ and suppose that $p_d(\mathbf{x}) = p_m(\mathbf{x}; \theta_*)$ as implied by the assumptions in Section 4.

It is too difficult to analytically derive the optimal noise distribution in general. To alleviate this problem, we make the following assumptions:

(B1) The noise distribution p_n is contained in $\mathcal{P}_{\epsilon_0}$ defined as follows: for sufficiently small positive constant $\epsilon_0 \ll 1$,

$$\mathcal{P}_{\epsilon_0} = \left\{ p : \Omega \rightarrow [0, \infty) \mid \int_{\Omega} p(\mathbf{x}) d\mu(\mathbf{x}) = 1, \right. \\ \left. \exists h, \exists \epsilon \in (0, \epsilon_0), \frac{p_d(\mathbf{x})}{p(\mathbf{x})} = 1 + \epsilon h(\mathbf{x}), \|h\|_{\infty} \leq 1 \right\}.$$

(B2) φ belongs to $\mathcal{C}^3$.

(B3) $p_n(\mathbf{y})$ is normalized so that $\int p_n(\mathbf{y}) d\mathbf{y} = 1$.

(B4) Fisher information matrix $\mathbf{I}_F$ is positive definite.

Assumption (B1) is striking, and enables us to have a first-order expression of MSE with respect to ϵ whose details are given in Appendix I. The following theorem presents the optimal noise distribution as the minimizer of the first-order expression:

Theorem 3 (Optimal noise distribution). *Under Assumptions (B1-4), the optimal noise distribution minimizing the first-order approximation of MSE is given by*

$$p_n^{\text{opt}}(\mathbf{y}) = \frac{p_d(\mathbf{y}) \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{y})\|}{Z_n}, \quad (18)$$

where $Z_n := \mathbb{E}_d[\|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{X})\|]$ is the normalizing constant.

The proof is given in Appendix I. The optimal noise distribution p_n^{opt} is exactly the same as the one derived in Chehab et al. [2022] and is clearly different from the data distribution p_d . As depicted in Figure 1, it seems to emphasize the region where p_d strongly changes. Interestingly, $p_n^{\text{opt}}(\mathbf{y})$ is optimal for *any* choice of convex functions φ .

It should be noted that there is a gap in the optimal noise distributions between Proposition 2 and Theorem 3: Theorem 3 shows that the optimal distribution is p_n^{opt} , while

Proposition 2 indicates that the data distribution p_d is optimal in the limit of $\nu \rightarrow \infty$. In order to bridge p_d and p_n^{opt} , we consider to take the following constraint into account in the optimization problem:

$$\int \frac{(p_d(\mathbf{x}) - p_n(\mathbf{x}))^2}{p_n(\mathbf{x})} d\mathbf{x} \leq C_{\epsilon}, \quad (19)$$

where C_{ϵ} denotes a non-negative constant. The left-hand side on (19) is the Pearson χ^2 -divergence, and thus (19) means that p_n is optimized in a neighborhood of p_d more explicitly than Theorem 3 as well as Chehab et al. [2022]. The optimal noise distribution is established under this constraint in the following theorem:

Theorem 4 (Optimal noise distribution in a neighborhood of p_d). *Assume that $(1 + C_{\epsilon})/Z_n \leq \int p_d(\mathbf{x}) / \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\| d\mathbf{x}$ holds. Then, the optimal noise distribution minimizing the first-order approximation of MSE under constraint (19) is given as follows: For $0 < \lambda < \infty$,*

$$\check{p}_n^{\text{opt}}(\mathbf{y}; \lambda) \propto \sqrt{p_n^{\text{opt}}(\mathbf{y})^2 + \lambda p_d(\mathbf{y})^2}, \quad (20)$$

where λ is determined by C_{ϵ} and satisfies (47). When $\lambda = \infty$, $\check{p}_n^{\text{opt}}(\mathbf{y}; \lambda) = p_d(\mathbf{y})$.

The proof provided in Appendix J is substantially different from that of Chehab et al. [2022], where a novel lemma is established. Interestingly, $\check{p}_n^{\text{opt}}(\mathbf{y}; \lambda)$ is a combination of $p_n^{\text{opt}}(\mathbf{y})$ and $p_d(\mathbf{y})$. If $\lambda = 0$, then $\check{p}_n^{\text{opt}}$ reduces to p_n^{opt} , while it approaches p_d as λ increases. This result might give some insight to the gap between the noise distribution in Chehab et al. [2022] and the common practical belief that p_d is the optimal noise distribution mentioned in Section 1.

Finally, we propose two constraints for robustness and derive the optimal noise distributions under them. Theorem 3 and Figure 1 imply that p_n^{opt} has a wider support than p_d , and thus distribution ratio $r_m(\mathbf{x}; \alpha_*) = p_d(\mathbf{x}) / p_n^{\text{opt}}(\mathbf{x}) \propto \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|^{-1}$ can approach zero when outliers $\mathbf{x}_o$ go to the tails of p_d . Based on this implication, the condition (12) for B-robustness can be simplified and generalized as follows: For all $\mathbf{x}$,

$$\frac{p_d(\mathbf{x})}{p_n(\mathbf{x})} \|\mathbf{g}_m(\mathbf{x})\|^{\gamma} \leq C', \quad (21)$$

where $C' > 0$ and $\gamma > 0$ are a positive constant and a tuning parameter, respectively. A larger exponent γ increases the divergent speed. Therefore, if (21) holds even for larger γ , BCE could be more robust. Motivated by this idea, we propose to add (21) as a constraint in optimizing p_n . Furthermore, inspired by (21), we propose another constraint of an integral form, which yields a smoother optimal noise distribution, as follows:

$$\int \frac{p_d(\mathbf{x})}{p_n(\mathbf{x})} \|\mathbf{g}_m(\mathbf{x})\|^{\gamma} d\mathbf{x} \leq C'', \quad (22)$$

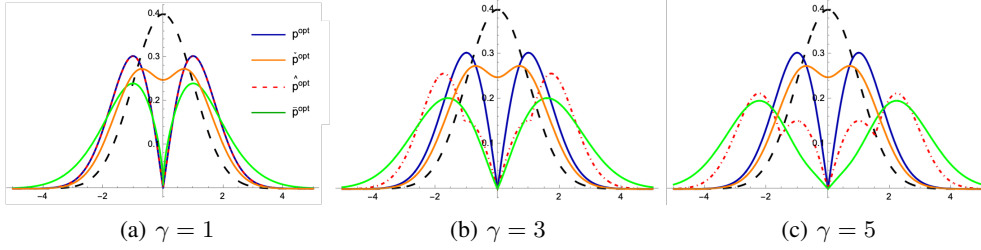


Figure 1: $p_n^{\text{opt}}(\mathbf{y})$, $\tilde{p}_n^{\text{opt}}(\mathbf{y}; \lambda)$, $\hat{p}_n^{\text{opt}}(\mathbf{y}; \gamma)$ and $\bar{p}_n^{\text{opt}}(\mathbf{y}; \lambda, \gamma)$ in mean estimation of the Gaussian distribution (black dotted line) when $\lambda = 0.5$.

where C'' is a positive constant. The following theorem 5 establishes the optimal noise distribution under each of the robust constraints:

Theorem 5 (Optimal noise distribution with robustness). *Under an additional constraint (21), the optimal noise distribution is given*

$$\hat{p}_n^{\text{opt}}(\mathbf{y}; \gamma) = \max(p_d(\mathbf{y}) \|\mathbf{g}_m(\mathbf{y})\|^\gamma / C', \lambda p_n^{\text{opt}}(\mathbf{y})), \quad (23)$$

where λ is chosen so that $\int \hat{p}_n^{\text{opt}}(\mathbf{y}; \lambda, \gamma) d\mathbf{y} = 1$.

Next, we replace constraint (21) by (22) and assume that $\int \sqrt{p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma} d\mathbf{x} \leq \sqrt{C''}$. Then, if $0 < \lambda < \infty$, the optimal noise distribution $\bar{p}_n^{\text{opt}}$ is given by

$$\bar{p}_n^{\text{opt}}(\mathbf{x}; \lambda, \gamma) \propto \sqrt{p_n^{\text{opt}}(\mathbf{y})^2 + \lambda p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma}, \quad (24)$$

where λ is determined by C'' and satisfies (49). When $\lambda = \infty$, $\bar{p}_n^{\text{opt}}(\mathbf{x}; \gamma) \propto \sqrt{p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma}$.

The proof is given in Appendix K and includes another lemma which has not been seen in Chehab et al. [2022]. Compared with the optimal noise distribution in Theorem 3, $\hat{p}_n^{\text{opt}}$ and $\bar{p}_n^{\text{opt}}$ in (24) have a wider support than p_n^{opt} . This wide support is important in robust estimation because it prevents the ratio $p_d(\mathbf{x})/p_n(\mathbf{x})$ from diverging when outliers go to the tails of $p_d(\mathbf{x})$. Indeed, Figure 1 illustrates $\hat{p}_n^{\text{opt}}$ and $\bar{p}_n^{\text{opt}}$ have a wider support than p_n^{opt} in Gaussian mean estimation.

6 NUMERICAL EXPERIMENTS

This section numerically investigates the robustness of τ -NCE and ρ -NCE without model selection, and demonstrates the performance of BCE with the optimal noise distribution derived under a robustness constraint. The experimental results with model selection based on cross-validation are given in Appendix L.

6.1 ROBUSTNESS OF BCE

We first generated the contaminated data from

$$(1 - \varepsilon)p_d(\mathbf{x}) + \varepsilon\bar{\delta}(\mathbf{x}), \quad (25)$$

where $\bar{\delta}(\mathbf{x})$ denotes a probability distribution of outliers. For p_d , we employed three kinds of distributions: Gaussian, log-normal, and Poisson distributions. Distributions in the same family were chosen for $\bar{\delta}(\mathbf{x})$ and $p_n(\mathbf{x})$. With unnormalized models, we estimated the ten-dimensional mean parameter vector for Gaussian and log-normal distributions, and the single parameter for Poisson distribution. The performance measure is RMSE to the true parameters. The details of the experiments are presented in Appendix L.

The results are presented in Fig. 2 when $\varepsilon = 0, 0.01, 0.05, 0.1$. For all data, both τ -NCE and ρ -NCE are robust against outliers, and compare favorably with PNCE. On the other hand, the estimation error of NCE quickly increases as the contamination ratio ε increases. These results support our theoretical analysis for B-robustness in Section 4.1 and possibly imply that BCE is robust under a variety of convex functions φ . As implied in Corollary 1, Figure 3 illustrates that the robustness is often improved as ν increases.

6.2 OPTIMAL NOISE DISTRIBUTION

Next, we investigate whether the optimal noise distribution $\bar{p}_n^{\text{opt}}(\mathbf{x}; \lambda, \gamma)$ in (23) produces robust parameter estimation. To this end, we generated one-dimensional data by the contaminated model (25) where $N(0, 1)$ and $N(\mu_o, 0.1^2)$ were employed for p_d and $\bar{\delta}$, respectively. The parameter μ_o in $N(\mu_o, 0.1^2)$ can be interpreted as the center of outliers. For comparison, p_d and p_n^{opt} were used as noise distribution p_n . Noise data samples were generated using the inverse function method for p_n^{opt} and slice sampling for $\bar{p}_n^{\text{opt}}$, respectively. As in Section 6.1, we estimated the mean value in the Gaussian distribution $N(0, 1)$ by PNCE with $\beta = 1$.

Figure 4 shows estimation errors when the contamination ratio ε is from 0.01 to 0.05. The optimal noise distribution $\bar{p}_n^{\text{opt}}(\mathbf{x}; \lambda, \gamma)$ produces very robust estimates, while the estimation by p_d and p_n^{opt} is strongly hampered by the more contamination of outliers. This robustness of $\bar{p}_n^{\text{opt}}(\mathbf{x}; \lambda, \gamma)$ is more clearly emphasized when μ_o is located around the heavier-tails of the data distribution p_d (Figure 4(b)). Estimation errors over the noise sample ratio $\hat{\nu}$ are also presented

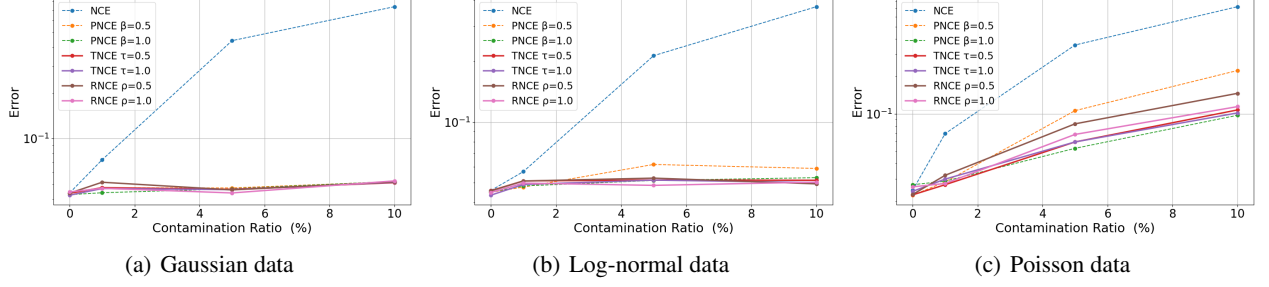


Figure 2: Estimation errors over outlier ratio ε when $(n, n_y) = (1000, 1000)$. Each point indicates the averaged errors over 100 runs, and the vertical axis is plotted on a logarithm scale.

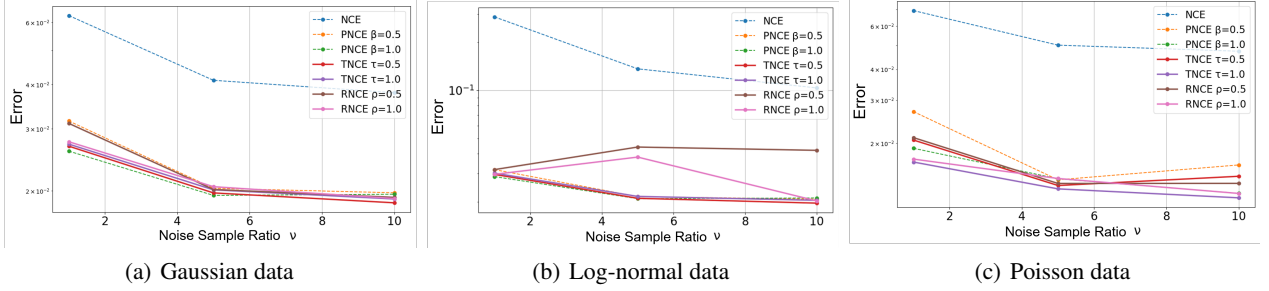


Figure 3: Estimation errors over sample ratio $\hat{\nu} = n_y/n$ when $(n, \varepsilon) = (3000, 0.01)$.

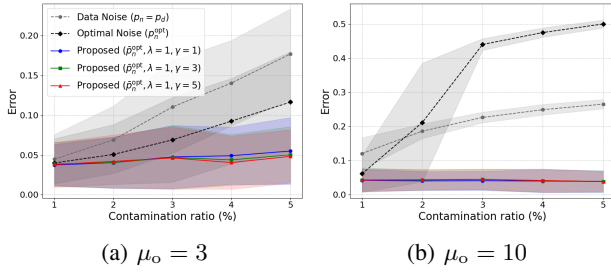


Figure 4: Estimation errors over contamination ratio ε when $(n, n_y) = (500, 2500)$. Each point is the average of estimator errors over 100 runs. The shaded regions represent the standard errors.

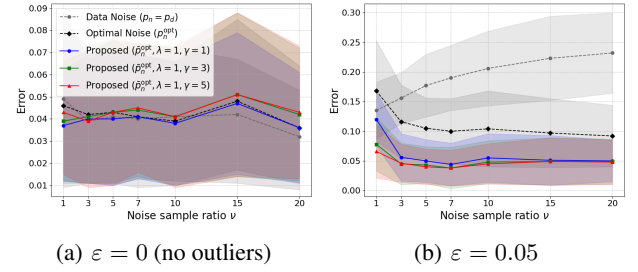


Figure 5: Estimation errors over sample ratio $\hat{\nu}$ when $(\mu_0, n) = (3, 500)$.

in Figure 5.

7 CONCLUSION

This paper proposed a general framework for learning unnormalized models based on the Bregman divergence. Our framework includes existing methods as special cases, and enables us to derive novel variants of NCE. We performed asymptotic analysis of our framework in terms of outlier-robustness and efficiency. Our analysis indicates that robust estimation is possible under a variety of convex functions in the Bregman divergence. The proposed framework can

achieve the Cramér-Rao bound when the noise distribution is equal to the data one. Inspired by the asymptotic analysis, the optimal noise distribution is also derived under novel constraints for robustness. We numerically demonstrated the robustness of the novel variants of NCE and our optimal noise distribution.

Acknowledgements

This work was partially supported by JSPS KAKENHI Grant Numbers JP23K28150, JP24K14849 and JP26H02491.

References

- J. Besag. Statistical analysis of non-lattice data. *Journal of the Royal Statistical Society*, 24D(3):179–195, 1975.
- L. M. Bregman. The relaxation method of finding the common point of convex sets and its application to the solution of problems in convex programming. *USSR Computational Mathematics and Mathematical Physics*, 7(3): 200–217, 1967.
- C. Ceylan and M. U. Gutmann. Conditional noise-contrastive estimation of unnormalised models. In *Proceedings of the 35th International Conference on Machine Learning (ICML)*, volume 80, pages 726–734. PMLR, 2018.
- O. Chehab, A. Gramfort, and A. Hyvärinen. The optimal noise in noise-contrastive learning is not what you think. In *Proceedings of the Thirty-Eighth Conference on Uncertainty in Artificial Intelligence (UAI)*, pages 307–316. PMLR, 2022.
- H. Chen, G. He, L. Yuan, G. Cui, H. Su, and J. Zhu. Noise contrastive alignment of language models with explicit rewards. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 37, pages 117784–117812, 2024.
- J. Feng, H. Xu, S. Mannor, and S. Yan. Robust logistic regression and classification. *Advances in neural information processing systems (NeurIPS)*, 27, 2014.
- R. Gao, E. Nijkamp, D. P. Kingma, Z. Xu, A. M. Dai, and Y. N. Wu. Flow contrastive estimation of energy-based models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020.
- I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio. Generative adversarial nets. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 2672–2680, 2014.
- M. Gutmann and A. Hyvärinen. Noise-contrastive estimation: A new estimation principle for unnormalized statistical models. In *Proceedings of the thirteenth international conference on artificial intelligence and statistics (AISTATS)*, pages 297–304. JMLR Workshop and Conference Proceedings, 2010.
- M. U. Gutmann and J. Hirayama. Bregman divergence as general framework to estimate unnormalized statistical models. In *Proceedings of the Twenty-Seventh Conference on Uncertainty in Artificial Intelligence (UAI)*, pages 283–290, 2011.
- M. U. Gutmann, S. Kleinegesse, and B. Rhodes. Statistical applications of contrastive learning. *Behaviormetrika*, 49 (2):277–301, 2022.
- M. U. Gutmann and A. Hyvärinen. Noise-contrastive estimation of unnormalized statistical models, with applications to natural image statistics. *Journal of Machine Learning Research*, 13:307–361, 2012.
- F. R. Hampel, E. M. Ronchetti, P. J. Rousseeuw, and W. A. Stahel. *Robust statistics: the approach based on influence functions*. John Wiley & Sons, 2011.
- G.E. Hinton. Training products of experts by minimizing contrastive divergence. *Neural Computation*, 14(8):1771–1800, 2002.
- P. J. Huber and E. M. Ronchetti. *Robust statistics*. Wiley, 2009.
- A. Hyvärinen. Estimation of non-normalized statistical models by score matching. *Journal of Machine Learning Research*, 6:695–709, 2005.
- T. Kanamori and M. Sugiyama. Statistical analysis of distance estimators with density differences and density ratios. *Entropy*, 16(2):921–942, 2014.
- P. Lavergne. A Cauchy-Schwarz inequality for expectation of matrices. *Discussion papers, Department of Economics, Simon Fraser University*, 2008.
- E. L. Lehmann and G. Casella. *Theory of Point Estimation*. Springer, 1998.
- X. Liu, F. Zhang, Z. Hou, Z. Wang, L. Mian, J. Zhang, and J. Tang. Self-supervised learning: Generative or contrastive. *arXiv preprint arXiv:2006.08218*, 2020.
- R. A. Maronna, R. D. Martin, V. J. Yohai, and M. Salibián-Barrera. *Robust statistics: theory and methods (with R)*. John Wiley & Sons, 2019.
- T. Matsuda, M. Uehara, and A. Hyvärinen. Information criteria for non-normalized models. *Journal of Machine Learning Research*, 22:1–33, 2021.
- M. Minami and S. Eguchi. Adaptive selection for minimum β -divergence method. In *Proceedings of the fourth international Symposium on Independent Component Analysis and Blind Source Separation*, 2003.
- A. Mnih and Y. W. Teh. A fast and simple algorithm for training neural probabilistic language models. In *Proceedings of the 29th International Conference on Machine Learning (ICML)*, 2012.
- N. Murata, . Takenouchi, T. Kanamori, and S. Eguchi. Information geometry of U-Boost and Bregman divergence. *Neural Computation*, 16(7):1437–1481, 2004.
- H. Sasaki and T. Takenouchi. Outlier-robust parameter estimation for unnormalized statistical models. *Japanese Journal of Statistics and Data Science*, 2024.

- Y. Song and S. Ermon. Generative modeling by estimating gradients of the data distribution. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 32, 2019.
- Y. Song, J. Sohl-Dickstein, D. P. Kingma, A. Kumar, S. Ermon, and B. Poole. Score-based generative modeling through stochastic differential equations. In *International Conference on Learning Representations (ICLR)*, 2021.
- T. Takenouchi and S. Eguchi. Robustifying AdaBoost by adding the naive error rate. *Neural Computation*, 16(4): 767–787, 2004.
- T. Takenouchi and T. Kanamori. Statistical inference with unnormalized discrete models and localized homogeneous divergences. *Journal of Machine Learning Research*, 18(1):1804–1829, 2017.
- M. Uehara, T. Kanamori, T. Takenouchi, and T. Matsuda. A unified statistically efficient estimation framework for unnormalized models. In *International Conference on Artificial Intelligence and Statistics (AISTATS)*, pages 809–819. PMLR, 2020.
- A. W. Van der Vaart. *Asymptotic statistics*. Cambridge university press, 1998.
- L. Wasserman. *All of statistics*. Springer, 2004.

Bregman contrastive estimation as a general framework for learning unnormalized models: Efficiency, robustness and optimal noise (Supplementary Material)

Yuto Fujii¹

Hiroaki Sasaki²

Takafumi Kanamori^{1,3}

¹Dept. of Math. & Comp. Sci., Institute of Science Tokyo, Tokyo, Japan

²Dept. of Math. Info., Meiji Gakuin University, Kanagawa, Japan

³RIKEN, Center for Advanced Intelligence Project, Tokyo, Japan

A NOTATION

Table 2: List of notations.

$\mathbf{x}, \mathbf{y}$	Input and noise data
$\ \mathbf{x}\ $	Euclidean norm of a vector $\mathbf{x}$
∇_{α}	Gradient with respect to α
n, n_y	Numbers of input and noise data samples
$p_d(\mathbf{x})$	Data distribution
$\bar{p}_d$	Contaminated distribution of p_d
$p_n(\mathbf{y})$	Noise distribution
$p_n^{\text{opt}}(\mathbf{y}), \check{p}_n^{\text{opt}}(\mathbf{y}; \lambda), \hat{p}_n^{\text{opt}}(\mathbf{y}; \gamma),$ $\bar{p}_n^{\text{opt}}(\mathbf{y}; \lambda, \gamma)$	Optimal noise distributions
$p_m(\mathbf{x}; \theta)$	Statistical model with a parameter vector θ
$p_m^0(\mathbf{x}; \alpha)$	Unnormalized statistical model with a parameter vector $\alpha = (\theta, c)$
$\mathcal{D} := \text{supp}(p_m)$	Support of p_m : $\{\mathbf{x} \mid p_m(\mathbf{x}; \alpha) \neq 0\}$
$\mathbf{g}_m(\mathbf{x}; \alpha)$	Gradient of $\log p_m^0(\mathbf{x}; \alpha)$ with respect to α : $\nabla_{\alpha} \log p_m^0(\mathbf{x}; \alpha)$
$\mathbf{g}_m(\mathbf{x})$	Gradient of $\log p_m(\mathbf{x}; \theta)$ with respect to θ at $\theta = \theta_*$: $\nabla_{\theta} \log p_m(\mathbf{x}; \theta) _{\theta=\theta_*}$
$r(\mathbf{x})$	Ratio of $p_d(\mathbf{x})$ to $p_n(\mathbf{x})$: $p_d(\mathbf{x})/p_n(\mathbf{x})$
$r_m(\mathbf{x}; \alpha)$	Ratio of $p_m^0(\mathbf{x}; \alpha)$ to $p_n(\mathbf{x})$: $p_m^0(\mathbf{x}; \alpha)/p_n(\mathbf{x})$
$f_C(\mathbf{x}; \alpha)$ ($C = 0, 1$)	Model for the posterior probability
ν	Ratio of class probabilities: $p(C = 1)/p(C = 0)$
$\hat{\nu}$	Sample ratio: n_y/n
$\varphi, \varphi', \varphi''$	Strictly convex function in the Bregman divergence and its derivatives
$\mathbf{x}_o$	Outlier
ε	Contamination ratio of outliers
$\delta_{\mathbf{x}_o}$	Dirac delta function with the point mass $\mathbf{x}_o$
$\mathbb{E}_d, \bar{\mathbb{E}}_d, \mathbb{E}_n$	Expectations over $p_d, \bar{p}_d$ and p_n
$\mathcal{L}_{\varphi}^{\text{SB}}(\alpha), \bar{\mathcal{L}}_{\varphi}^{\text{SB}}(\alpha)$	Risk function of BCE and its contaminated risk using $\bar{p}_d$
$\alpha_*, \alpha_{\varepsilon}$	Minimizers of $\mathcal{L}_{\varphi}^{\text{SB}}(\alpha)$ and $\bar{\mathcal{L}}_{\varphi}^{\text{SB}}(\alpha)$
$\Lambda_{\varphi}, \Lambda_{\text{NCE}}$	Asymptotic variances of BCE and NCE
$\mathbf{I}_F$	Fisher information matrix
$f_C^*(u)$	Model with the optimal parameter: $f_C(\mathbf{x}; \alpha_*)$
$\mathcal{C}^k$	Space of continuous and k -times differentiable functions

B DERIVATION OF (4)

Here, we first express $p(\mathbf{u}) = p_0\{p_d(\mathbf{u}) + \nu p_n(\mathbf{u})\}$ where $p_0 := p(C = 0)$. Then, (3) can be written as

$$\begin{aligned}
& - \int p(\mathbf{u}) \sum_{C \in \{0,1\}} [\varphi(f_C(\mathbf{u}; \boldsymbol{\alpha})) - \varphi'(f_C(\mathbf{u}; \boldsymbol{\alpha}))(p(C|\mathbf{u}) - f_C(\mathbf{u}; \boldsymbol{\alpha}))] d\mathbf{u} \\
& = p_0 \int (p_d(\mathbf{u}) + \nu p_n(\mathbf{u})) \sum_{C \in \{0,1\}} [-\varphi(f_C(\mathbf{u}; \boldsymbol{\alpha})) + \varphi'(f_C(\mathbf{u}; \boldsymbol{\alpha}))f_C(\mathbf{u}; \boldsymbol{\alpha})] d\mathbf{u} \\
& \quad - \sum_{C \in \{0,1\}} p(C) \int p(\mathbf{u}|C) \varphi'(f_C(\mathbf{u}; \boldsymbol{\alpha})) d\mathbf{u} \\
& = p_0 \left(\mathbb{E}_d \left[\sum_{C \in \{0,1\}} \{-\varphi(f_C(\mathbf{X}; \boldsymbol{\alpha})) + \varphi'(f_C(\mathbf{X}; \boldsymbol{\alpha}))f_C(\mathbf{X}; \boldsymbol{\alpha})\} \right] \right. \\
& \quad \left. + \nu \mathbb{E}_n \left[\sum_{C \in \{0,1\}} \{-\varphi(f_C(\mathbf{Y}; \boldsymbol{\alpha})) + \varphi'(f_C(\mathbf{Y}; \boldsymbol{\alpha}))f_C(\mathbf{Y}; \boldsymbol{\alpha})\} \right] - \mathbb{E}_d[\varphi'(f_0(\mathbf{X}; \boldsymbol{\alpha}))] - \nu \mathbb{E}_n[\varphi'(f_1(\mathbf{Y}; \boldsymbol{\alpha}))] \right) \\
& = p_0 \left(-\mathbb{E}_d \left[\varphi(f_0(\mathbf{X}; \boldsymbol{\alpha})) + \varphi(f_1(\mathbf{X}; \boldsymbol{\alpha})) - \varphi'(f_0(\mathbf{X}; \boldsymbol{\alpha}))\{f_0(\mathbf{X}; \boldsymbol{\alpha}) - 1\} - \varphi'(f_1(\mathbf{X}; \boldsymbol{\alpha}))f_1(\mathbf{X}; \boldsymbol{\alpha}) \right] \right. \\
& \quad \left. - \nu \mathbb{E}_n \left[\varphi(f_0(\mathbf{Y}; \boldsymbol{\alpha})) + \varphi(f_1(\mathbf{Y}; \boldsymbol{\alpha})) - \varphi'(f_0(\mathbf{Y}; \boldsymbol{\alpha}))f_0(\mathbf{Y}; \boldsymbol{\alpha}) - \varphi'(f_1(\mathbf{Y}; \boldsymbol{\alpha}))\{f_1(\mathbf{Y}; \boldsymbol{\alpha}) - 1\} \right] \right).
\end{aligned}$$

Since $f_0(\mathbf{u}; \boldsymbol{\alpha}) + f_1(\mathbf{u}; \boldsymbol{\alpha}) = 1$ for all $\mathbf{u}$, we finally obtain

$$\begin{aligned}
\mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}) & = -\mathbb{E}_d[\varphi(f_0(\mathbf{X}; \boldsymbol{\alpha})) + \varphi(f_1(\mathbf{X}; \boldsymbol{\alpha})) + \{\varphi'(f_0(\mathbf{X}; \boldsymbol{\alpha})) - \varphi'(f_1(\mathbf{X}; \boldsymbol{\alpha}))\}f_1(\mathbf{X}; \boldsymbol{\alpha})] \\
& \quad - \nu \mathbb{E}_n[\varphi(f_0(\mathbf{Y}; \boldsymbol{\alpha})) + \varphi(f_1(\mathbf{Y}; \boldsymbol{\alpha})) - \{\varphi'(f_0(\mathbf{Y}; \boldsymbol{\alpha})) - \varphi'(f_1(\mathbf{Y}; \boldsymbol{\alpha}))\}f_0(\mathbf{Y}; \boldsymbol{\alpha})].
\end{aligned}$$

C PROOF OF PROPOSITION 1

It follows from the definitions of $\boldsymbol{\alpha}_\star$ and $\boldsymbol{\alpha}_\varepsilon$ that

$$\nabla_\alpha \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha})|_{\boldsymbol{\alpha}=\boldsymbol{\alpha}_\star} = \mathbf{0} \quad \text{and} \quad \nabla_\alpha \bar{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha})|_{\boldsymbol{\alpha}=\boldsymbol{\alpha}_\varepsilon} = \mathbf{0}. \tag{26}$$

A simple calculation shows that $\nabla_\alpha \bar{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha})$ satisfies

$$\nabla_\alpha \bar{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}) = \nabla_\alpha \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}) - \varepsilon \{ \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha})f_1(\mathbf{X}; \boldsymbol{\alpha})g_m(\mathbf{X}; \boldsymbol{\alpha})] - w_\varphi(\mathbf{x}_o; \boldsymbol{\alpha})f_1(\mathbf{x}_o; \boldsymbol{\alpha})g_m(\mathbf{x}_o; \boldsymbol{\alpha}) \},$$

where we employed the following formulae:

$$\nabla_\alpha f_0(\mathbf{x}; \boldsymbol{\alpha}) = f_0(\mathbf{x}; \boldsymbol{\alpha})f_1(\mathbf{x}; \boldsymbol{\alpha})g_m(\mathbf{x}; \boldsymbol{\alpha}) \quad \text{and} \quad \nabla_\alpha f_1(\mathbf{x}; \boldsymbol{\alpha}) = -f_0(\mathbf{x}; \boldsymbol{\alpha})f_1(\mathbf{x}; \boldsymbol{\alpha})g_m(\mathbf{x}; \boldsymbol{\alpha}). \tag{27}$$

Since the influence function is defined by a Gâteaux derivative of $\boldsymbol{\alpha}_\varepsilon$, differentiating the second equation in (26) with respect to ε at $\varepsilon = 0$ yields

$$\mathbf{H}_\varphi \cdot \text{IF}(\mathbf{x}_o) - \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_\star)f_1(\mathbf{X}; \boldsymbol{\alpha}_\star)g_m(\mathbf{X}; \boldsymbol{\alpha}_\star)] + w_\varphi(\mathbf{x}_o; \boldsymbol{\alpha}_\star)f_1(\mathbf{x}_o; \boldsymbol{\alpha}_\star)g_m(\mathbf{x}_o; \boldsymbol{\alpha}_\star) = \mathbf{0},$$

where $\mathbf{H}_\varphi := \nabla_\alpha^2 \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha})|_{\boldsymbol{\alpha}=\boldsymbol{\alpha}_\star}$. The invertibility assumption of $\mathbf{H}_\varphi$ completes the proof.

D EXAMPLES OF THE WEIGHT FUNCTION w_φ

This appendix gives examples of the weight function w_φ and illustrates them in Figure 6.

- PNCE:

$$w_\varphi(\mathbf{x}; \boldsymbol{\alpha}) = (f_0(\mathbf{x}; \boldsymbol{\alpha})^{\beta-1} + f_1(\mathbf{x}; \boldsymbol{\alpha})^{\beta-1}) f_0(\mathbf{x}; \boldsymbol{\alpha})f_1(\mathbf{x}; \boldsymbol{\alpha}).$$

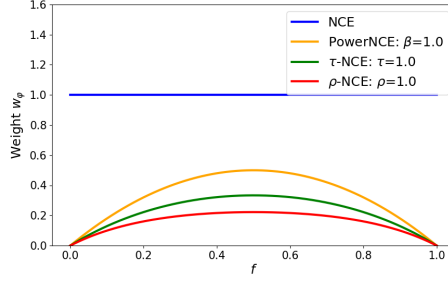


Figure 6: Illustration for the weight function $\{\varphi''(f) + \varphi''(1-f)\}f(1-f)$.

- τ -NCE:

$$w_\varphi(\mathbf{x}; \boldsymbol{\alpha}) = \left(\frac{1}{f_0(\mathbf{x}; \boldsymbol{\alpha}) + \tau} + \frac{1}{f_1(\mathbf{x}; \boldsymbol{\alpha}) + \tau} \right) f_0(\mathbf{x}; \boldsymbol{\alpha}) f_1(\mathbf{x}; \boldsymbol{\alpha}).$$

- ρ -NCE:

$$w_\varphi(\mathbf{x}; \boldsymbol{\alpha}) = \left(\frac{1}{(f_0(\mathbf{x}; \boldsymbol{\alpha}) + \rho)^2} + \frac{1}{(f_1(\mathbf{x}; \boldsymbol{\alpha}) + \rho)^2} \right) f_0(\mathbf{x}; \boldsymbol{\alpha}) f_1(\mathbf{x}; \boldsymbol{\alpha}).$$

E CONSISTENCY OF BCE

Proposition 3 (Consistency). *Suppose that Assumptions (A1-6) hold. Then, $\hat{\boldsymbol{\alpha}}$ converges to $\boldsymbol{\alpha}_*$ in probability.*

Proof. We essentially follow the consistency proof in Van der Vaart [1998]. By definition of $\boldsymbol{\alpha}_*$ and Assumptions (A3-4), we have $\boldsymbol{\theta}_0 = \boldsymbol{\theta}_*$. Thus, it suffices that $\hat{\boldsymbol{\alpha}}$ converges $\boldsymbol{\alpha}_*$ in probability.

From Assumption (A5), an event $\{\|\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}_*\| \geq \epsilon\}$ for each $\epsilon > 0$ implies another event $\{\mathcal{L}_\varphi^{\text{SB}}(\hat{\boldsymbol{\alpha}}) > \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) + \delta\}$ for some $\delta > 0$. Thus, we have the following inequality of these events:

$$P(\|\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}_*\| \geq \epsilon) \leq P(\mathcal{L}_\varphi^{\text{SB}}(\hat{\boldsymbol{\alpha}}) - \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) > \delta).$$

Since $\hat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}) > \hat{\mathcal{L}}_\varphi^{\text{SB}}(\hat{\boldsymbol{\alpha}})$ holds for all $\boldsymbol{\alpha}$ by definition of $\hat{\boldsymbol{\alpha}}$,

$$\begin{aligned} \mathcal{L}_\varphi^{\text{SB}}(\hat{\boldsymbol{\alpha}}) - \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) &= \mathcal{L}_\varphi^{\text{SB}}(\hat{\boldsymbol{\alpha}}) - \hat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) + \hat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) - \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) \\ &\leq \mathcal{L}_\varphi^{\text{SB}}(\hat{\boldsymbol{\alpha}}) - \hat{\mathcal{L}}_\varphi^{\text{SB}}(\hat{\boldsymbol{\alpha}}) + \hat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) - \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) \\ &\leq 2 \sup_{\boldsymbol{\alpha}} |\mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}) - \hat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha})|. \end{aligned}$$

Therefore, Assumption (A6) ensures that $\mathcal{L}_\varphi^{\text{SB}}(\hat{\boldsymbol{\alpha}}) - \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*)$ converges to zero in probability, indicating the convergence of $\hat{\boldsymbol{\alpha}}$ to $\boldsymbol{\alpha}_*$. The proof is completed. $\square$

F PROOF OF THEOREM 1

By definition of $\hat{\boldsymbol{\alpha}}$, we have

$$\nabla_{\boldsymbol{\alpha}} \hat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha})|_{\boldsymbol{\alpha}=\hat{\boldsymbol{\alpha}}} = -\frac{1}{n} \sum_{i=1}^n w_\varphi(\mathbf{x}_i; \hat{\boldsymbol{\alpha}}) f_1(\mathbf{x}_i; \hat{\boldsymbol{\alpha}}) \mathbf{g}_m(\mathbf{x}_i; \hat{\boldsymbol{\alpha}}) + \frac{\hat{\nu}}{n_y} \sum_{i=1}^{n_y} w_\varphi(\mathbf{y}_i; \hat{\boldsymbol{\alpha}}) f_0(\mathbf{y}_i; \hat{\boldsymbol{\alpha}}) \mathbf{g}_m(\mathbf{y}_i; \hat{\boldsymbol{\alpha}}) = \mathbf{0}. \quad (28)$$

The Taylor expansion of (28) around $\boldsymbol{\alpha}_*$ yields

$$\nabla_{\boldsymbol{\alpha}} \mathcal{L}_\varphi^{\text{SB}}(\boldsymbol{\alpha})|_{\boldsymbol{\alpha}=\boldsymbol{\alpha}_*} + \hat{\mathbf{H}}_\varphi(\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}_*) + O(\|\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}_*\|^2) = \mathbf{0}, \quad (29)$$

where

$$\begin{aligned}
\widehat{\mathbf{H}}_\varphi &:= \nabla_{\boldsymbol{\alpha}}^2 \widehat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha})|_{\boldsymbol{\alpha}=\boldsymbol{\alpha}_*} \\
&= -\frac{1}{n} \sum_{i=1}^n [\{\nabla_{\boldsymbol{\alpha}} w_\varphi(\mathbf{x}_i; \boldsymbol{\alpha}_*)\} f_1^*(\mathbf{x}_i) \mathbf{g}_m(\mathbf{x}_i; \boldsymbol{\alpha}_*)^\top + w_\varphi(\mathbf{x}_i; \boldsymbol{\alpha}_*) \{\nabla_{\boldsymbol{\alpha}} f_1^*(\mathbf{x}_i)\} \mathbf{g}_m(\mathbf{x}_i; \boldsymbol{\alpha}_*)^\top \\
&\quad + w_\varphi(\mathbf{x}_i; \boldsymbol{\alpha}_*) f_1^*(\mathbf{x}_i) \nabla_{\boldsymbol{\alpha}} \mathbf{g}_m(\mathbf{x}_i; \boldsymbol{\alpha}_*)] + \frac{\hat{\nu}}{n_y} \sum_{i=1}^{n_y} [\{\nabla_{\boldsymbol{\alpha}} w_\varphi(\mathbf{y}_i; \boldsymbol{\alpha}_*)\} f_0^*(\mathbf{y}_i) \mathbf{g}_m(\mathbf{y}_i; \boldsymbol{\alpha}_*)^\top \\
&\quad + w_\varphi(\mathbf{y}_i; \boldsymbol{\alpha}_*) \{\nabla_{\boldsymbol{\alpha}} f_0^*(\mathbf{y}_i)\} \mathbf{g}_m(\mathbf{y}_i; \boldsymbol{\alpha}_*)^\top + w_\varphi(\mathbf{y}_i; \boldsymbol{\alpha}_*) f_0^*(\mathbf{y}_i) \nabla_{\boldsymbol{\alpha}} \mathbf{g}_m(\mathbf{y}_i; \boldsymbol{\alpha}_*)],
\end{aligned} \tag{30}$$

where $f_0^*(\mathbf{x}) := 1 - f_1^*(\mathbf{x}) = \frac{p_d(\mathbf{x})}{p_d(\mathbf{x}) + \nu p_n(\mathbf{x})}$ and $\nabla_{\boldsymbol{\alpha}} f_1^*(\mathbf{x}) := \nabla_{\boldsymbol{\alpha}} f_1(\mathbf{x}; \boldsymbol{\alpha})|_{\boldsymbol{\alpha}=\boldsymbol{\alpha}_*}$. Since $\nu p_n f_0^* = p_d f_1^*$, the expectation of $\widehat{\mathbf{H}}_\varphi$ over $\mathcal{X}$ and $\mathcal{Y}$ is computed as

$$\begin{aligned}
\mathbb{E}_{\mathcal{X} \times \mathcal{Y}}[\widehat{\mathbf{H}}_\varphi] &= \int w_\varphi(\mathbf{x}; \boldsymbol{\alpha}_*) f_0^*(\mathbf{x}) f_1^*(\mathbf{x}) \mathbf{g}_m(\mathbf{x}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{x}; \boldsymbol{\alpha}_*)^\top p_d(\mathbf{x}) d\mathbf{x} \\
&\quad + \int w_\varphi(\mathbf{y}; \boldsymbol{\alpha}_*) f_0^*(\mathbf{y}) f_1^*(\mathbf{y}) \mathbf{g}_m(\mathbf{y}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{y}; \boldsymbol{\alpha}_*)^\top \hat{\nu} p_n(\mathbf{y}) d\mathbf{y} \\
&= \int \left\{ f_0^*(\mathbf{x}) + \frac{\hat{\nu}}{\nu} f_1^*(\mathbf{x}) \right\} w_\varphi(\mathbf{x}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{x}) \mathbf{g}_m(\mathbf{x}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{x}; \boldsymbol{\alpha}_*)^\top p_d(\mathbf{x}) d\mathbf{x} \\
&= \mathbb{E}_d \left[\left\{ f_0^*(\mathbf{X}) + \frac{\hat{\nu}}{\nu} f_1^*(\mathbf{X}) \right\} w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top \right].
\end{aligned}$$

Since $\hat{\nu}$ is assumed to converge to ν , $f_0^*(\mathbf{x}) + \frac{\hat{\nu}}{\nu} f_1^*(\mathbf{x}) \rightarrow 1$ as $n \rightarrow \infty$ and $n_y \rightarrow \infty$. Therefore, the law of large numbers guarantees that $\nabla_{\boldsymbol{\alpha}} \widehat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*)$ converges to $\mathbf{J}_\varphi = \mathbb{E}_d [w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top]$ in probability.

On the other hand, by definition of $\boldsymbol{\alpha}_*$, the expectation of $\nabla_{\boldsymbol{\alpha}} \widehat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha})$ at $\boldsymbol{\alpha} = \boldsymbol{\alpha}_*$ satisfies

$$\mathbb{E}_{\mathcal{X} \times \mathcal{Y}}[\nabla_{\boldsymbol{\alpha}} \widehat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*)] = \mathbf{0}, \tag{31}$$

The variance of $\nabla_{\boldsymbol{\alpha}} \widehat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha})$ at $\boldsymbol{\alpha} = \boldsymbol{\alpha}_*$ can be computed as follows:

$$\begin{aligned}
&\text{Var}_{\mathcal{X} \times \mathcal{Y}}[\nabla_{\boldsymbol{\alpha}} \widehat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*)] \\
&= \mathbb{E}_{\mathcal{X} \times \mathcal{Y}} \left[\left(-\frac{1}{n} \sum_{i=1}^n w_\varphi(\mathbf{X}_i; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}_i) \mathbf{g}_m(\mathbf{X}_i; \boldsymbol{\alpha}_*) + \frac{\hat{\nu}}{n_y} \sum_{i=1}^{n_y} w_\varphi(\mathbf{Y}_i; \boldsymbol{\alpha}_*) f_0^*(\mathbf{Y}_i) \mathbf{g}_m(\mathbf{Y}_i; \boldsymbol{\alpha}_*) \right) \right. \\
&\quad \left. \left(-\frac{1}{n} \sum_{i=1}^n w_\varphi(\mathbf{X}_i; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}_i) \mathbf{g}_m(\mathbf{X}_i; \boldsymbol{\alpha}_*) + \frac{\hat{\nu}}{n_y} \sum_{i=1}^{n_y} w_\varphi(\mathbf{Y}_i; \boldsymbol{\alpha}_*) f_0^*(\mathbf{Y}_i) \mathbf{g}_m(\mathbf{Y}_i; \boldsymbol{\alpha}_*) \right)^\top \right] \\
&= \frac{1}{n^2} \mathbb{E}_{\mathcal{X}} \left[\left(\sum_{i=1}^n w_\varphi(\mathbf{X}_i; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}_i) \mathbf{g}_m(\mathbf{X}_i; \boldsymbol{\alpha}_*) \right) \left(\sum_{i=1}^n w_\varphi(\mathbf{X}_i; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}_i) \mathbf{g}_m(\mathbf{X}_i; \boldsymbol{\alpha}_*) \right)^\top \right] \\
&\quad + \frac{\hat{\nu}^2}{n_y^2} \mathbb{E}_{\mathcal{Y}} \left[\left(\sum_{i=1}^{n_y} w_\varphi(\mathbf{Y}_i; \boldsymbol{\alpha}_*) f_0^*(\mathbf{Y}_i) \mathbf{g}_m(\mathbf{Y}_i; \boldsymbol{\alpha}_*) \right) \left(\sum_{i=1}^{n_y} w_\varphi(\mathbf{Y}_i; \boldsymbol{\alpha}_*) f_0^*(\mathbf{Y}_i) \mathbf{g}_m(\mathbf{Y}_i; \boldsymbol{\alpha}_*) \right)^\top \right] \\
&\quad - 2\hat{\nu} \mathbb{E}_d [w_\varphi(\mathbf{X}) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)] \mathbb{E}_n [w_\varphi(\mathbf{Y}) f_0^*(\mathbf{Y}) \mathbf{g}_m(\mathbf{Y}; \boldsymbol{\alpha}_*)]^\top \\
&= \frac{1}{n} \mathbb{E}_d [w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top] \\
&\quad - \frac{1}{n} \left(1 + \frac{1}{\hat{\nu}} \right) \mathbb{E}_d [w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)] \mathbb{E}_d [w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)]^\top
\end{aligned}$$

Hence, it follows from the central limit theorem that $\sqrt{n} \nabla_{\boldsymbol{\alpha}} \widehat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*)$ converges to a normal distribution as follows:

$$\sqrt{n} \nabla_{\boldsymbol{\alpha}} \widehat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*) \xrightarrow{D} N(\mathbf{0}, \mathbf{I}_\varphi), \tag{32}$$

where

$$\begin{aligned} \mathbf{I}_\varphi &:= \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top] \\ &\quad - \left(1 + \frac{1}{\nu}\right) \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)] \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)]^\top. \end{aligned}$$

By the dominant convergence theorem and the consistency of $\hat{\boldsymbol{\alpha}}$, higher-order terms in (29) can be ignored, which yields

$$\sqrt{n}(\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}_*) = \sqrt{n} \mathbf{H}_\varphi^{-1} \nabla_{\boldsymbol{\alpha}} \hat{\mathcal{L}}_\varphi^{\text{SB}}(\boldsymbol{\alpha}_*).$$

Thus, Slutsky's theorem [Van der Vaart, 1998] and (32) ensures that $\sqrt{n}(\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}_*)$ converges to a normal distribution as follows:

$$\sqrt{n}(\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}_*) \xrightarrow{D} N(\mathbf{0}, \mathbf{J}_\varphi^{-1} \mathbf{I}_\varphi \mathbf{J}_\varphi^{-1}).$$

This completes the proof.

G PROOF OF THEOREM 2

As derived in Gutmann and Hyvärinen [2012], the asymptotic variance of NCE is given by

$$\boldsymbol{\Lambda}_{\text{NCE}} := \mathbf{I}^{-1} \left(\mathbf{I} - \left(1 + \frac{1}{\nu}\right) \mathbf{m} \mathbf{m}^\top \right) \mathbf{I}^{-1}, \quad (33)$$

where

$$\mathbf{I} := \mathbb{E}_d[f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top] \quad \text{and} \quad \mathbf{m} := \mathbb{E}_d[f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)].$$

By definition of $\mathbf{g}_m(\mathbf{u}; \boldsymbol{\alpha})$, we have

$$\mathbf{g}_m(\mathbf{u}; \boldsymbol{\alpha}) = \nabla_{\boldsymbol{\alpha}} \log p_m^0(\mathbf{u}; \boldsymbol{\alpha}) = (\nabla_{\boldsymbol{\theta}} \log p_m(\mathbf{u}; \boldsymbol{\theta})^\top, -1)^\top. \quad (34)$$

Then, by letting a canonical vector be $\mathbf{e}_1 := (0, \dots, 0, 1)$, it holds $\mathbf{g}_m(\mathbf{u}; \boldsymbol{\alpha}_*)^\top \mathbf{e}_1 = -1$, based on which the second term on the right-hand side of $\mathbf{I}_\varphi$ in (14) can be computed as

$$\begin{aligned} &\mathbb{E}_d[f_1^*(\mathbf{X}) w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)] \mathbb{E}_d[f_1^*(\mathbf{X}) w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)]^\top \\ &= \mathbb{E}_d[f_1^*(\mathbf{X}) w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top \mathbf{e}_1] \mathbb{E}_d[f_1^*(\mathbf{X}) w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top \mathbf{e}_1]^\top \\ &= \mathbb{E}_d[f_1^*(\mathbf{X}) w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top] \mathbf{e}_1 \mathbf{e}_1^\top \mathbb{E}_d[f_1^*(\mathbf{X}) w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top] \\ &= \mathbf{J}_\varphi \mathbf{e}_1 \mathbf{e}_1^\top \mathbf{J}_\varphi. \end{aligned} \quad (35)$$

Eq.(35) enables us to express $\boldsymbol{\Lambda}_\varphi$ as

$$\boldsymbol{\Lambda}_\varphi = \mathbf{J}_\varphi^{-1} \mathbf{I}_\varphi \mathbf{J}_\varphi^{-1} = \mathbf{J}_\varphi^{-1} (\mathbb{E}_d[f_1^*(\mathbf{x}) w_\varphi^2(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top]) \mathbf{J}_\varphi^{-1} - \left(1 + \frac{1}{\nu}\right) \mathbf{e}_1 \mathbf{e}_1^\top. \quad (36)$$

Applying a matrix version of the Cauchy-Schwarz inequality [Lavergne, 2008] to the first term in (36) gives

$$\mathbf{J}_\varphi^{-1} (\mathbb{E}_d[f_1^*(\mathbf{X}) w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*)^2 \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top]) \mathbf{J}_\varphi^{-1} \succeq \mathbb{E}_d[f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)^\top]^{-1} = \mathbf{I}^{-1}. \quad (37)$$

On the other hand, the second term in (36) can be computed as

$$\mathbf{e}_1 \mathbf{e}_1^\top = \mathbf{I}^{-1} \{ \mathbf{I} (\mathbf{e}_1 \mathbf{e}_1^\top) \mathbf{I} \} \mathbf{I}^{-1} = \mathbf{I}^{-1} (\mathbf{m} \mathbf{m}^\top) \mathbf{I}^{-1},$$

where we applied the following equation derived similarly as (35):

$$\mathbf{m} \mathbf{m}^\top = \mathbb{E}_d[f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)] \mathbb{E}_d[f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \boldsymbol{\alpha}_*)]^\top = \mathbf{I} \mathbf{e}_1 \mathbf{e}_1^\top \mathbf{I}. \quad (38)$$

Finally, by combining (36) with (37) and (38), we obtain

$$\boldsymbol{\Lambda}_\varphi \succeq \mathbf{I}^{-1} - \left(1 + \frac{1}{\nu}\right) \mathbf{I}^{-1} (\mathbf{m} \mathbf{m}^\top) \mathbf{I}^{-1} = \mathbf{I}^{-1} \left\{ \mathbf{I} - \left(1 + \frac{1}{\nu}\right) \mathbf{m} \mathbf{m}^\top \right\} \mathbf{I}^{-1} = \boldsymbol{\Lambda}_{\text{NCE}}.$$

Equality holds if the following equation holds almost everywhere: For all $\mathbf{u}$, there exist constant vectors $\mathbf{a}, \mathbf{b}$ such that

$$w_\varphi(\mathbf{x}; \boldsymbol{\alpha}_*) \sqrt{f_1^*(\mathbf{x})} \mathbf{a}^\top \mathbf{g}_m(\mathbf{x}; \boldsymbol{\alpha}_*) + \sqrt{f_1^*(\mathbf{x})} \mathbf{b}^\top \mathbf{g}_m(\mathbf{x}; \boldsymbol{\alpha}_*) = 0,$$

implying that $w_\varphi(\mathbf{u}; \boldsymbol{\alpha}_*)$ is constant by assuming that $\|\mathbf{g}_m(\mathbf{u}; \boldsymbol{\alpha}_*)\|$ and $\sqrt{f_1^*(\mathbf{u})}$ are positive.

H PROOF OF PROPOSITION 2

It follows from the definition of α_* and Assumptions (A3-4) that $p_d(\cdot) = p_m^0(\cdot; \alpha_*)$ holds, giving

$$f_0^*(\mathbf{u}) = \frac{p_m^0(\mathbf{u}; \alpha_*)}{p_m^0(\mathbf{u}; \alpha_*) + \nu p_n(\mathbf{u})} = \frac{p_d(\mathbf{u})}{p_d(\mathbf{u}) + \nu p_n(\mathbf{u})} = \frac{p_d(\mathbf{u})}{\nu p_n(\mathbf{u})} + O(\nu^{-2}) \quad (39)$$

$$f_1^*(\mathbf{u}) = \frac{\nu p_n(\mathbf{u})}{p_m^0(\mathbf{u}; \alpha_*) + \nu p_n(\mathbf{u})} = \frac{\nu p_n(\mathbf{u})}{p_d(\mathbf{u}) + \nu p_n(\mathbf{u})} = 1 - \frac{p_d(\mathbf{u})}{\nu p_n(\mathbf{u})} + O(\nu^{-2}), \quad (40)$$

By assumption that φ is of class $\mathcal{C}^2$, φ'' is bounded on $[0, 1]$. Thus, we can define the limit values as $C_0 = \lim_{z \rightarrow 0} \varphi''(z)$ and $C_1 = \lim_{z \rightarrow 1} \varphi''(z)$. Then, the weight function w_φ has the following order expression:

$$w_\varphi(\mathbf{u}; \alpha_*) = \frac{(C_0 + C_1)p_d(\mathbf{u})}{\nu p_n(\mathbf{u})} + O(\nu^{-2}). \quad (41)$$

Substituting (40) and (41) into $\mathbf{J}_\varphi$ gives

$$\mathbf{J}_\varphi = (C_0 + C_1)\mathbb{E}_d \left[\frac{p_d(\mathbf{X})}{\nu p_n(\mathbf{X})} \mathbf{g}_m(\mathbf{X}; \alpha_*) \mathbf{g}_m(\mathbf{X}; \alpha_*)^\top \right] + O(\nu^{-2}) = \frac{C_0 + C_1}{\nu} \tilde{\mathbf{J}} + O(\nu^{-2}). \quad (42)$$

Similarly, regarding the first term in $\mathbf{I}_\varphi$, we have

$$\mathbb{E}_d[w_\varphi(\mathbf{X}; \alpha_*)^2 f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \alpha_*) \mathbf{g}_m(\mathbf{X}; \alpha_*)^\top] = \frac{(C_0 + C_1)^2}{\nu^2} \mathbb{E}_d \left[\left(\frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} \right)^2 \mathbf{g}_m(\mathbf{X}; \alpha_*) \mathbf{g}_m(\mathbf{X}; \alpha_*)^\top \right] + O(\nu^{-3}),$$

while the second term is written as

$$\begin{aligned} & \left(1 + \frac{1}{\nu} \right) \mathbb{E}_d[w_\varphi(\mathbf{X}; \alpha_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \alpha_*)] \mathbb{E}_d[w_\varphi(\mathbf{X}; \alpha_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}; \alpha_*)] \\ &= \frac{(C_0 + C_1)^2}{\nu^2} \mathbb{E}_d \left[\left(\frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} \right) \mathbf{g}_m(\mathbf{X}; \alpha_*) \right] \mathbb{E}_d \left[\left(\frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} \right) \mathbf{g}_m(\mathbf{X}; \alpha_*) \right]^\top + O(\nu^{-3}). \end{aligned}$$

By using these expressions in $\mathbf{I}_\varphi$, we obtain

$$\mathbf{I}_\varphi = \frac{(C_0 + C_1)^2}{\nu^2} \tilde{\mathbf{I}} + O(\nu^{-3}). \quad (43)$$

Since

$$\mathbf{J}_\varphi^{-1} = \left\{ \frac{C_0 + C_1}{\nu} \right\}^{-1} \tilde{\mathbf{J}}^{-1} + O(\nu^{-2}),$$

we have $\mathbf{J}_\varphi^{-1} \mathbf{I}_\varphi \mathbf{J}_\varphi^{-1} = \tilde{\mathbf{J}}^{-1} \tilde{\mathbf{I}} \tilde{\mathbf{J}}^{-1} + O(\nu^{-2})$, which completes the proof.

I PROOF OF THEOREM 3

Since $p_d(\mathbf{x}) = p_m(\mathbf{x}; \theta_*)$, $f_1^*(\mathbf{x})$ can be represented as in (13). Then, $\frac{p_d(\mathbf{x})}{p_n(\mathbf{x})} = 1 + \epsilon h(\mathbf{x})$ in Assumption (B1) gives the following order expression:

$$\begin{aligned} f_1^*(\mathbf{x}) &= \frac{\nu p_n(\mathbf{x})}{p_d(\mathbf{x}) + \nu p_n(\mathbf{x})} = \frac{1}{1 + \frac{1}{\nu} + \frac{\epsilon h(\mathbf{x})}{\nu}} = \frac{\nu}{\nu + 1} - \frac{\nu}{(\nu + 1)^2} \epsilon h(\mathbf{x}) + O(\epsilon^2) \\ f_0^*(\mathbf{x}) &= \frac{p_d(\mathbf{x})}{p_d(\mathbf{x}) + \nu p_n(\mathbf{x})} = \frac{\frac{1}{\nu} + \epsilon \frac{h(\mathbf{x})}{\nu}}{1 + \frac{1}{\nu} + \frac{\epsilon h(\mathbf{x})}{\nu}} = \frac{1}{\nu + 1} + \frac{\nu}{(\nu + 1)^2} \epsilon h(\mathbf{x}) + O(\epsilon^2). \end{aligned}$$

These order expressions enable us to represent $w_\varphi(\mathbf{x}; \boldsymbol{\alpha}_*)$ as

$$\begin{aligned}
w_\varphi(\mathbf{x}; \boldsymbol{\alpha}_*) &= \{\varphi''(f_0^*(\mathbf{x})) + \varphi''(f_1^*(\mathbf{x}))\} f_0^*(\mathbf{x}) f_1^*(\mathbf{x}) \\
&= \left\{ \varphi''\left(\frac{1}{\nu+1}\right) + \varphi''' \left(\frac{1}{\nu+1}\right) \frac{\nu}{(\nu+1)^2} \epsilon h(\mathbf{x}) + \varphi''\left(\frac{\nu}{\nu+1}\right) - \varphi''' \left(\frac{\nu}{\nu+1}\right) \frac{\nu}{(\nu+1)^2} \epsilon h(\mathbf{x}) + O(\epsilon^2) \right\} \\
&\quad \times \left\{ \frac{\nu}{(\nu+1)^2} + \frac{\nu^2 - \nu}{(\nu+1)^3} \epsilon h(\mathbf{x}) + O(\epsilon^2) \right\} \\
&= C_1 + C_2 \epsilon h(\mathbf{x}) + O(\epsilon^2),
\end{aligned}$$

where

$$\begin{aligned}
C_1 &:= \left\{ \varphi''\left(\frac{\nu}{\nu+1}\right) + \varphi''\left(\frac{1}{\nu+1}\right) \right\} \frac{\nu}{(\nu+1)^2} \\
C_2 &:= \left\{ \varphi''\left(\frac{\nu}{\nu+1}\right) + \varphi''\left(\frac{1}{\nu+1}\right) \right\} \frac{\nu^2 - \nu}{(\nu+1)^3} + \left\{ \varphi''' \left(\frac{1}{\nu+1}\right) - \varphi''' \left(\frac{\nu}{\nu+1}\right) \right\} \frac{\nu^2}{(\nu+1)^4}.
\end{aligned}$$

In order to compute $\mathbf{I}_\varphi$ and $\mathbf{J}_\varphi$, we substitute this expression of $w_\varphi(\mathbf{x}; \boldsymbol{\alpha}_*)$ into

$$\begin{aligned}
w_\varphi(\mathbf{x}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{x}) &= \{C_1 + C_2 \epsilon h(\mathbf{x}) + O(\epsilon^2)\} \left\{ \frac{\nu}{\nu+1} - \frac{\nu}{(\nu+1)^2} \epsilon h(\mathbf{x}) + O(\epsilon^2) \right\} \\
&= \frac{C_1 \nu}{\nu+1} + \left\{ \frac{C_2 \nu}{\nu+1} - \frac{C_1 \nu}{(\nu+1)^2} \right\} \epsilon h(\mathbf{x}) + O(\epsilon^2) \\
w_\varphi(\mathbf{x}; \boldsymbol{\alpha}_*)^2 f_1^*(\mathbf{x}) &= \{C_1^2 + 2C_1 C_2 \epsilon h(\mathbf{x}) + O(\epsilon^2)\} \left\{ \frac{\nu}{\nu+1} - \frac{\nu}{(\nu+1)^2} \epsilon h(\mathbf{x}) + O(\epsilon^2) \right\} \\
&= \frac{C_1^2 \nu}{\nu+1} + \left\{ \frac{2C_1 C_2 \nu}{\nu+1} - \frac{C_1^2 \nu}{(\nu+1)^2} \right\} \epsilon h(\mathbf{x}) + O(\epsilon^2).
\end{aligned}$$

Then, $\mathbf{I}_\varphi$ and $\mathbf{J}_\varphi$ are expressed as

$$\begin{aligned}
\mathbf{I}_\varphi &= \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*)^2 f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}) \mathbf{g}_m(\mathbf{X})^\top] \\
&\quad - \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X})] \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X})]^\top \\
&= C_1^2 \frac{\nu}{\nu+1} \mathbf{I}_F + \left\{ \frac{2C_1 C_2 \nu}{\nu+1} - \frac{C_1^2 \nu}{(\nu+1)^2} \right\} \mathbf{A}(\epsilon) + O(\epsilon^2) \\
\mathbf{J}_\varphi &= \mathbb{E}_d[w_\varphi(\mathbf{X}; \boldsymbol{\alpha}_*) f_1^*(\mathbf{X}) \mathbf{g}_m(\mathbf{X}) \mathbf{g}_m(\mathbf{X})^\top] = C_1 \frac{\nu}{\nu+1} \mathbf{I}_F + \left\{ \frac{C_2 \nu}{\nu+1} - \frac{C_1 \nu}{(\nu+1)^2} \right\} \mathbf{A}(\epsilon) + O(\epsilon^2),
\end{aligned}$$

where $\mathbf{A}(\epsilon) := \epsilon \mathbb{E}_d[h(\mathbf{X}) \mathbf{g}_m(\mathbf{X}) \mathbf{g}_m(\mathbf{X})^\top]$ and we applied $\mathbb{E}_d[\mathbf{g}_m(\mathbf{X})] = 0$. Since

$$\begin{aligned}
\mathbf{J}_\varphi^{-1} &= \left\{ C_1 \frac{\nu}{\nu+1} \mathbf{I}_F + \left\{ \frac{C_2 \nu}{\nu+1} - \frac{C_1 \nu}{(\nu+1)^2} \right\} \mathbf{A}(\epsilon) + O(\epsilon^2) \right\}^{-1} \\
&= \frac{\nu+1}{C_1 \nu} \mathbf{I}_F^{-1} - \left(\frac{C_2(\nu+1)}{C_1^2 \nu} - \frac{1}{C_1 \nu} \right) \mathbf{I}_F^{-1} \mathbf{A}(\epsilon) \mathbf{I}_F^{-1} + O(\epsilon^2),
\end{aligned}$$

we compute the asymptotic variance $\boldsymbol{\Lambda}_\varphi$ as follows:

$$\boldsymbol{\Lambda}_\varphi = \mathbf{J}_\varphi^{-1} \mathbf{I}_\varphi \mathbf{J}_\varphi^{-1} = \frac{\nu+1}{\nu} \mathbf{I}_F^{-1} + \frac{1}{\nu} \mathbf{I}_F^{-1} \mathbf{A}(\epsilon) \mathbf{I}_F^{-1} + O(\epsilon^2).$$

The first-order approximation of $\text{MSE} = \text{Tr}(\boldsymbol{\Lambda}_\varphi)/n$ is given by

$$\text{MSE} = \frac{\nu+1}{n\nu} \text{Tr}(\mathbf{I}_F^{-1}) + \frac{1}{n\nu} \text{Tr}(\mathbf{I}_F^{-1} \mathbf{A}(\epsilon) \mathbf{I}_F^{-1}) + O(\epsilon^2) \quad (44)$$

$$= \frac{\nu+1}{n\nu} \text{Tr}(\mathbf{I}_F^{-1}) + \frac{\epsilon}{n\nu} \mathbb{E}_d \left[\left\{ \frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} - 1 \right\} \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{X})\|^2 \right], \quad (45)$$

where since $\mathbf{I}_F$ is symmetric and $\epsilon h(\mathbf{x}) = \frac{p_d(\mathbf{x})}{p_n(\mathbf{x})} - 1$, we computed the second term as

$$\text{Tr}(\mathbf{I}_F^{-1} \mathbf{A}(\epsilon) \mathbf{I}_F^{-1}) = \epsilon \mathbb{E}_d [\text{Tr}(\{\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{X})\} \{\mathbf{g}_m(\mathbf{X}) \mathbf{I}_F^{-1}\}^\top)] = \epsilon \mathbb{E}_d \left[\left\{ \frac{p_d(\mathbf{X})}{p_n(\mathbf{X})} - 1 \right\} \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{X})\|^2 \right].$$

Next, we derive the optimal noise distribution as the minimizer of the second term on the right-hand side of (45) with a constraint $\int p_n(\mathbf{x}) d\mathbf{x} = 1$. Hence, we optimize the following Lagrangian:

$$\min_{p_n} \left[\int \frac{p_d(\mathbf{x})^2}{p_n(\mathbf{x})} \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|^2 d\mathbf{x} + \lambda \int p_n(\mathbf{x}) d\mathbf{x} \right], \quad (46)$$

where λ denotes the Lagrange multiplier. By computing the variational derivative and setting it to zero, we have

$$p_n^{\text{opt}}(\mathbf{y}) \propto p_d(\mathbf{y}) \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{y})\|,$$

which completes the proof.

J PROOF OF THEOREM 4

Our proof relies on the following Lemma whose proof is provided below:

Lemma 1. *Let $q(\mathbf{x})$ be a probability density and let $u(\mathbf{x}) \geq 0$ be an integrable function. Consider the optimization problem*

$$\min_{p:\text{p.d.f}} \int \frac{q(\mathbf{x})^2}{p(\mathbf{x})} d\mathbf{x} \quad \text{such that} \quad \int \frac{u(\mathbf{x})^2}{p(\mathbf{x})} d\mathbf{x} \leq C,$$

where C is a positive constant. Then,

- (a) *If $C < (\int u(\mathbf{x}) d\mathbf{x})^2$, then the problem is infeasible.*
- (b) *If $(\int u(\mathbf{x}) d\mathbf{x})^2 \leq C < \int \frac{u(\mathbf{x})^2}{q(\mathbf{x})} d\mathbf{x}$, then there exists $\lambda_0 \in (0, \infty]$ such that*

$$\left(\int \frac{u(\mathbf{x})^2}{\sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}} d\mathbf{x} \right) \left(\int \sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2} d\mathbf{x} \right) = C,$$

and the optimal solution is

$$p_{\text{opt}}(\mathbf{x}) = \frac{\sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}}{\int \sqrt{q(\mathbf{z})^2 + \lambda_0 u(\mathbf{z})^2} d\mathbf{z}}.$$

With the convention $\lambda_0 = \infty$, which corresponds to $C = (\int u(\mathbf{x}) d\mathbf{x})^2$, we obtain

$$p_{\text{opt}}(\mathbf{x}) = \frac{u(\mathbf{x})}{\int u(\mathbf{z}) d\mathbf{z}}$$

- (c) *If $C \geq \int \frac{u(\mathbf{x})^2}{q(\mathbf{x})} d\mathbf{x}$, then $p_{\text{opt}}(\mathbf{x}) = q(\mathbf{x})$.*

We first rewrite the Pearson χ^2 divergence in (19) as

$$\int \frac{(p_d(\mathbf{x}) - p_n(\mathbf{x}))^2}{p_n(\mathbf{x})} d\mathbf{x} = \int \frac{p_d(\mathbf{x})^2}{p_n(\mathbf{x})} d\mathbf{x} - 1.$$

Then, our optimization problem is formulated as

$$\min_{p_n:\text{p.d.f}} \int \frac{p_d(\mathbf{x})^2}{p_n(\mathbf{x})} \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|^2 d\mathbf{x} \quad \text{such that} \quad \int \frac{p_d(\mathbf{x})^2}{p_n(\mathbf{x})} d\mathbf{x} \leq 1 + C_\epsilon.$$

To apply Lemma 1, we next substitute $q(\mathbf{x}) = p_d(\mathbf{x})\|\mathbf{I}_F^{-1}\mathbf{g}_m(\mathbf{x})\|/Z_n$ and $u(\mathbf{x}) = p_d(\mathbf{x})$. Since C_ϵ is a non-negative constant, it holds $(\int u(\mathbf{x})d\mathbf{x})^2 = 1 \leq 1 + C_\epsilon$ and thus, Case (a) in Lemma 1 is excluded. On the other hand, the left-hand side of the condition in Case (c) of Lemma 1 is given by

$$\int \frac{u(\mathbf{x})^2}{q(\mathbf{x})}d\mathbf{x} = Z_n \int \frac{p_d(\mathbf{x})}{\|\mathbf{I}_F^{-1}\mathbf{g}_m(\mathbf{x})\|}d\mathbf{x}.$$

The assumption excludes Case (c). Hence, the optimal noise distribution is Case (b) and given by

$$\bar{p}_n^{\text{opt}}(\mathbf{x}; \gamma, \lambda) \propto \begin{cases} \sqrt{p_d(\mathbf{x})^2\|\mathbf{I}_F^{-1}\mathbf{g}_m(\mathbf{x})\|^2/Z_n^2 + \lambda p_d(\mathbf{x})^2} & 0 < \lambda < \infty \\ \sqrt{p_d(\mathbf{x})\|\mathbf{g}_m(\mathbf{x})\|^{\gamma/2}} & \lambda = \infty, \end{cases}$$

where λ is chosen such that

$$1 + C_\epsilon = \left(\int \frac{p_d(\mathbf{x})^2}{\sqrt{p_d(\mathbf{x})^2\|\mathbf{I}_F^{-1}\mathbf{g}_m(\mathbf{x})\|^2/Z_n^2 + \lambda p_d(\mathbf{x})^2}}d\mathbf{x} \right) \left(\int \sqrt{p_d(\mathbf{x})^2\|\mathbf{I}_F^{-1}\mathbf{g}_m(\mathbf{x})\|^2/Z_n^2 + \lambda p_d(\mathbf{x})^2}d\mathbf{x} \right). \quad (47)$$

Proof of Lemma 1. (a) For any probability density p , by Cauchy-Schwarz inequality,

$$\left(\int u(\mathbf{x})d\mathbf{x} \right)^2 = \left(\int \frac{u(\mathbf{x})}{\sqrt{p(\mathbf{x})}} \sqrt{p(\mathbf{x})}d\mathbf{x} \right)^2 \leq \left(\int \frac{u(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x} \right) \left(\int p(\mathbf{x})d\mathbf{x} \right) = \int \frac{u(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x}.$$

Hence $\int u(\mathbf{x})^2/p(\mathbf{x})d\mathbf{x} \geq (\int u(\mathbf{x})d\mathbf{x})^2$, so the problem is infeasible when $C < (\int u(\mathbf{x})d\mathbf{x})^2$.

(c) For any probability density p ,

$$\int \frac{q(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x} = \int \frac{(q(\mathbf{x}) - p(\mathbf{x}))^2}{p(\mathbf{x})}d\mathbf{x} + 1 \geq 1,$$

with equality if and only if $p = q$ a.e. If $C \geq \int u(\mathbf{x})^2/q(\mathbf{x})d\mathbf{x}$, then $p = q$ is feasible and thus optimal.

(b) Define, for $\lambda > 0$,

$$h(\lambda) := \left(\int \frac{u(\mathbf{x})^2}{\sqrt{q(\mathbf{x})^2 + \lambda u(\mathbf{x})^2}}d\mathbf{x} \right) \left(\int \sqrt{q(\mathbf{x})^2 + \lambda u(\mathbf{x})^2}d\mathbf{x} \right).$$

The dominated convergence theorem guarantees that $h(\lambda)$ for $0 < \lambda < \infty$ is continuous. Indeed, the inequalities

$$\begin{aligned} \frac{u(\mathbf{x})^2}{\sqrt{q(\mathbf{x})^2 + \lambda u(\mathbf{x})^2}} &\leq \frac{u(\mathbf{x})}{\sqrt{\lambda}}, \\ \sqrt{q(\mathbf{x})^2 + \lambda u(\mathbf{x})^2} &\leq q(\mathbf{x}) + \sqrt{\lambda}u(\mathbf{x}) \end{aligned}$$

and the integrability of q and u lead to the continuity of $h(\lambda)$. Since $h(0) = \int u(\mathbf{x})^2/q(\mathbf{x})d\mathbf{x}$ and $\lim_{\lambda \rightarrow \infty} h(\lambda) = (\int u(\mathbf{x})d\mathbf{x})^2$, there exists $\lambda_0 \in (0, \infty]$ such that $h(\lambda_0) = C$ under the assumption of (b).

Assume first that $\lambda_0 < \infty$. Let p be any feasible probability density. Then,

$$\int \frac{q(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x} = \int \frac{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x} - \lambda_0 \int \frac{u(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x} \geq \int \frac{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x} - \lambda_0 C.$$

By Cauchy-Schwarz inequality,

$$\int \frac{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x} = \int \frac{(\sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2})^2}{p(\mathbf{x})}d\mathbf{x} \geq \frac{\left(\int \sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}d\mathbf{x} \right)^2}{\int p(\mathbf{x})d\mathbf{x}} = \left(\int \sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}d\mathbf{x} \right)^2.$$

Therefore,

$$\int \frac{q(\mathbf{x})^2}{p(\mathbf{x})}d\mathbf{x} \geq \left(\int \sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}d\mathbf{x} \right)^2 - \lambda_0 C. \quad (48)$$

Let $p_{\text{opt}}(\mathbf{x}) \propto \sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}$ by which the equality holds in (48). Then, the Cauchy–Schwarz inequality is tight, and

$$\int \frac{u(\mathbf{x})^2}{p_{\text{opt}}(\mathbf{x})} d\mathbf{x} = \left(\int \sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2} d\mathbf{x} \right) \left(\int \frac{u(\mathbf{x})^2}{\sqrt{q(\mathbf{x})^2 + \lambda_0 u(\mathbf{x})^2}} d\mathbf{x} \right) = h(\lambda_0) = C,$$

so p_{opt} is feasible and attains the lower bound. Hence, it is optimal.

Finally, when $\lambda_0 = \infty$, we have $(\int u(\mathbf{x}) d\mathbf{x})^2 = C$, which means that the constraint is

$$\int \frac{u(\mathbf{x})^2}{p(\mathbf{x})} d\mathbf{x} = \left(\int u(\mathbf{x}) d\mathbf{x} \right)^2 \int \frac{\bar{u}(\mathbf{x})^2}{p(\mathbf{x})} d\mathbf{x} = \left(\int u(\mathbf{x}) d\mathbf{x} \right)^2 \left(\int \frac{(\bar{u}(\mathbf{x}) - p(\mathbf{x}))^2}{p(\mathbf{x})} d\mathbf{x} + 1 \right)^2 \leq \left(\int u(\mathbf{x}) d\mathbf{x} \right)^2$$

for $\bar{u}(\mathbf{x}) = u(\mathbf{x}) / \int u(\mathbf{z}) d\mathbf{z}$. Thus, the only feasible solution is $p = \bar{u}$. $\square$

K PROOF OF THEOREM 5

K.1 ROBUST CONSTRAINT

Our optimization problem is formulated as

$$\min_{p_n: \text{p.d.f.}} \int \frac{p_d(\mathbf{x})^2}{p_n(\mathbf{x})} \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|^2 d\mathbf{x} \quad \text{such that} \quad \forall \mathbf{x}, \frac{p_d(\mathbf{x})}{p_n(\mathbf{x})} \|\mathbf{g}_m(\mathbf{x})\|^\gamma \leq C'.$$

The proof is completed by substituting $q(\mathbf{x}) = p_d(\mathbf{x}) \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\| / \mathbb{E}_d[\|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{X})\|]$ and $u(\mathbf{x}) = p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma / C'$ in the following Lemma:

Lemma 2. *Let $q(\mathbf{x})$ be a probability density and $u(\mathbf{x})$ a non-negative function. Consider the optimization problem*

$$\min_{p: \text{p.d.f.}} \int \frac{q(\mathbf{x})^2}{p(\mathbf{x})} d\mathbf{x} \quad \text{such that} \quad \forall \mathbf{x}, p(\mathbf{x}) \geq u(\mathbf{x}).$$

Suppose that $\int u(\mathbf{x}) d\mathbf{x} \leq 1$. Then the optimal solution is given by

$$p_{\text{opt}}(\mathbf{x}) = \max\{u(\mathbf{x}), c_0 q(\mathbf{x})\},$$

where $c_0 \geq 0$ is a non-negative constant chosen so that $\int p_{\text{opt}}(\mathbf{x}) d\mathbf{x} = 1$. Moreover we have $c_0 \in [0, 1]$.

Proof of Lemma 2. Define the function

$$c \mapsto \eta(c) := \int \max\{u(\mathbf{x}), cq(\mathbf{x})\} d\mathbf{x}.$$

The function $\eta(c)$ is increasing and satisfies $\eta(0) = \int u(\mathbf{x}) d\mathbf{x} \leq 1$ and $\eta(1) = \int \max\{u(\mathbf{x}), q(\mathbf{x})\} d\mathbf{x} \geq \int q(\mathbf{x}) d\mathbf{x} = 1$. Moreover, by the dominated convergence theorem, $\eta(c)$ is continuous on $[0, 1]$. Hence, by the intermediate value theorem, there exists $c_0 \in [0, 1]$ such that $\eta(c_0) = 1$. If $\eta(0) = 1$, then $p(\mathbf{x}) = u(\mathbf{x})$ (with $c_0 = 0$) is the unique feasible solution. In the below, we assume $c_0 > 0$.

For a probability density $p(\mathbf{x})$, define the functional

$$F(p) := \int \frac{q(\mathbf{x})^2}{p(\mathbf{x})} d\mathbf{x} = \int \left(\frac{q(\mathbf{x})^2}{p(\mathbf{x})} + \frac{p(\mathbf{x})}{c_0^2} \right) d\mathbf{x} - \frac{1}{c_0^2}.$$

For any probability density $p(\mathbf{x})$ satisfying $p(\mathbf{x}) \geq u(\mathbf{x})$, we have

$$\frac{q(\mathbf{x})^2}{p(\mathbf{x})} + \frac{p(\mathbf{x})}{c_0^2} \geq \frac{q(\mathbf{x})^2}{p_{\text{opt}}(\mathbf{x})} + \frac{p_{\text{opt}}(\mathbf{x})}{c_0^2}$$

for all $\mathbf{x}$. Therefore, $F(p) \geq F(p_{\text{opt}})$ for all feasible p , which completes the proof. $\square$

K.2 ROBUST-INTEGRAL CONSTRAINT

The proof is based on Lemma 1 and is similar as Theorem 4. We first formulate our optimization problem as

$$\min_{p_n: \text{p.d.f.}} \int \frac{p_d(\mathbf{x})^2}{p_n(\mathbf{x})} \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|^2 d\mathbf{x} \quad \text{such that} \quad \int \frac{p_d(\mathbf{x})}{p_n(\mathbf{x})} \|\mathbf{g}_m(\mathbf{x})\|^\gamma d\mathbf{x} \leq C''.$$

To apply Lemma 1, we next substitute $q(\mathbf{x}) = p_d(\mathbf{x}) \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|/Z_n$ and $u(\mathbf{x}) = \sqrt{p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma}$. By assumption, Case (a) in Lemma 1 is excluded. After the substitution, the left-hand side of the condition in Case (c) of Lemma 1 can be expressed as

$$\int \frac{u(\mathbf{x})^2}{q(\mathbf{x})} d\mathbf{x} = \frac{Z_n}{C''} \int \frac{\|\mathbf{g}_m(\mathbf{x})\|^\gamma}{\|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|} d\mathbf{x} \geq \frac{Z_n}{C'' \kappa_{\max}} \int \|\mathbf{g}_m(\mathbf{x})\|^{\gamma-1} d\mathbf{x},$$

where $\kappa_{\max}$ denotes the maximum eigenvalue of $\mathbf{I}_F$. Since the right-hand side above diverges in general, Case (c) is also excluded. Hence, the optimal noise distribution is Case (b) and given by

$$\bar{p}_n^{\text{opt}}(\mathbf{x}; \gamma, \lambda) \propto \begin{cases} \sqrt{p_d(\mathbf{x})^2 \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|^2 / Z_n^2 + \lambda p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma} & 0 < \lambda < \infty \\ \sqrt{p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma} & \lambda = \infty, \end{cases}$$

where λ is chosen such that

$$C'' = \left(\int \frac{p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma}{\sqrt{p_d(\mathbf{x})^2 \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|^2 / Z_n^2 + \lambda p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma}} d\mathbf{x} \right) \left(\int \sqrt{p_d(\mathbf{x})^2 \|\mathbf{I}_F^{-1} \mathbf{g}_m(\mathbf{x})\|^2 / Z_n^2 + \lambda p_d(\mathbf{x}) \|\mathbf{g}_m(\mathbf{x})\|^\gamma} d\mathbf{x} \right). \quad (49)$$

L EXPERIMENTAL DETAILS AND RESULTS WITH CROSS-VALIDATION

For $p_d(\mathbf{x})$, $\bar{\delta}(\mathbf{x})$ and $p_n(\mathbf{x})$, we employed the following probability distributions:

- **Gaussian data** ($d = 10$): The standard normal density $N(\mathbf{0}, \mathbf{I}_d)$ was used for $p_d(\mathbf{x})$, while $N(\mathbf{3}, 0.1^2 \mathbf{I}_d)$ was for $\bar{\delta}(\mathbf{x})$ where $N(\mathbf{m}, \mathbf{C})$ denotes the Gaussian density with the mean vector $\mathbf{m}$ and covariance matrix $\mathbf{C}$, and $\mathbf{I}_d$ is the d by d identity matrix. Then, we estimate the mean vector $\mathbf{0}$ in p_d with the following unnormalized model:

$$\log p_m^0(\mathbf{x}; \boldsymbol{\alpha}) = -\frac{1}{2} \sum_{i=1}^d (x_i - \theta_i)^2 - c.$$

As done in NCE [Gutmann and Hyvärinen, 2010], we used the Gaussian density with the sample mean and sample covariance for $p_n(\mathbf{x})$. Estimation errors were measured by $\|\hat{\boldsymbol{\theta}} - \mathbf{0}\|/\sqrt{d}$ where $\hat{\boldsymbol{\theta}}$ denotes an estimate of $(\theta_1, \dots, \theta_d)$.

- **Log-normal data** ($d = 10$): The following log-normal density was used for $p_d(\mathbf{x})$ and $\bar{\delta}(\mathbf{x})$:

$$\prod_{i=1}^d \frac{1}{x_i \sqrt{2\pi\sigma^2}} \exp\left(-\frac{(\log x_i - \mu)^2}{2\sigma^2}\right), \quad (50)$$

where $x_i > 0$ ($i = 1, \dots, d$), and we used $\mu = 1$ and $\mu = 3$ for $p_d(\mathbf{x})$ and $\bar{\delta}(\mathbf{x})$ respectively, and $\sigma = 1$ in both distributions. An unnormalized model,

$$\log p_m^0(\mathbf{x}; \boldsymbol{\alpha}) = -\sum_{i=1}^d \left\{ \log x_i + \frac{1}{2} (\log x_i - \theta_i)^2 \right\} - c,$$

was adopted to estimate the parameter $\mu = 1$ in $p_d(\mathbf{x})$, and estimation errors were measured by $\|\hat{\boldsymbol{\theta}} - \mathbf{1}\|/\sqrt{d}$ where $\mathbf{1} = (1, \dots, 1)^\top$. For $p_n(\mathbf{y})$, the log-normal density was employed where the empirical estimates were substituted into μ and σ in (50).

- **Poisson data** ($d = 1$): We sampled one-dimensional data x by using Poisson distribution:

$$\frac{\lambda^x e^{-\lambda}}{x!},$$

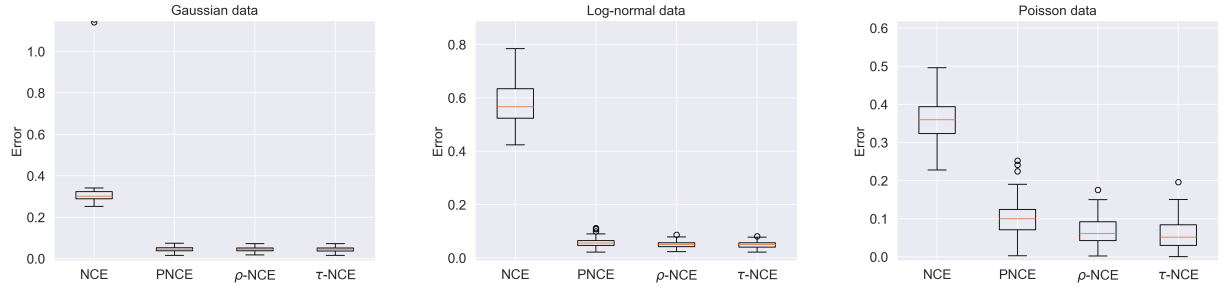
where λ is a positive parameter. For data generation, $\lambda = 0.5$ and $\lambda = 5$ were used for p_d and $\bar{\delta}(x)$, respectively. As done in Sasaki and Takenouchi [2024], we employed Poisson distribution with the parameter value twice larger than the empirical estimate of λ for p_n . To estimate λ , our unnormalized model is formulated as

$$\log p_m^0(x; \boldsymbol{\alpha}) = x \log \theta + \log x! - c.$$

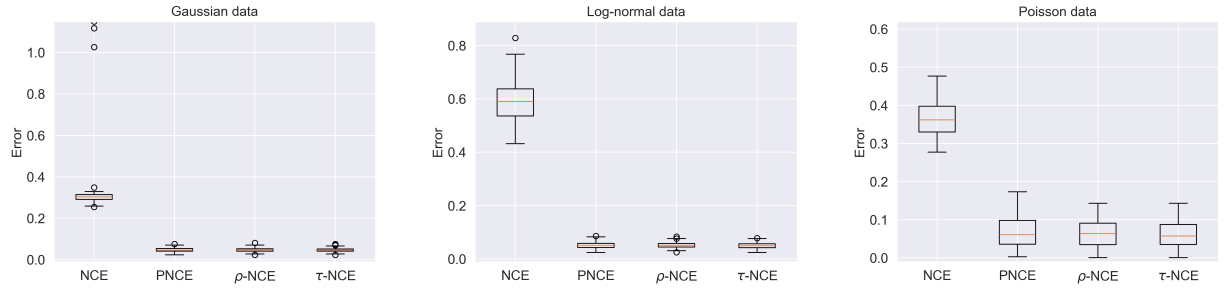
Estimation errors were measured by $|\hat{\theta} - 0.5|$.

All parameters of unnormalized models were optimized by BFGS.

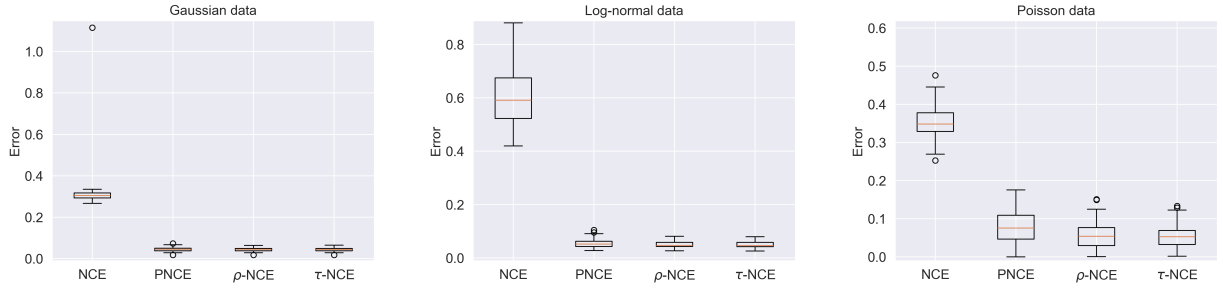
The same experiments as Section 6 were performed with cross-validation in Section 3.1. Figure 7 indicates that our cross-validation method fairly works well on a range of the anchor parameters.



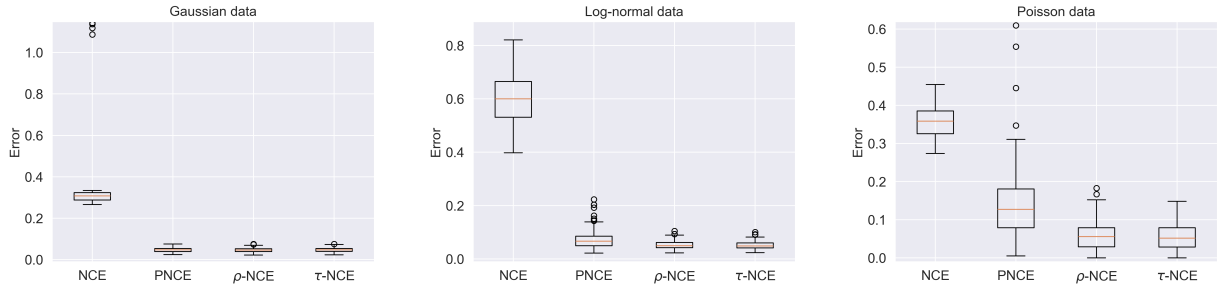
(a) $\beta_0 = \rho_0 = \tau_0 = 0.5$



(b) $\beta_0 = \rho_0 = \tau_0 = 1$



(c) $\beta_0 = \rho_0 = \tau_0 = 3$



(d) $\beta_0 = \rho_0 = \tau_0 = 5$

Figure 7: Estimation errors of NCE, PNCE, ρ -NCE and τ -NCE when $(n, n_y, \epsilon) = (1000, 1000, 0.05)$ and anchor parameters $(\beta_0, \rho_0$ and $\tau_0)$ are changed. The hyper-parameters of PNCE, ρ -NCE and τ -NCE are selected based on five-fold cross validation.