
Locally Optimal Private Sampling: Beyond the Global Minimax

Hrad Ghoukasian

Department of Computing and Software
McMaster University
ghoukash@mcmaster.ca

Bonwoo Lee

Department of Mathematical Sciences
Korea Advanced Institute of Science & Technology
righthim@kaist.ac.kr

Shahab Asodeh

Department of Computing and Software
McMaster University
asodeh@mcmaster.ca

Abstract

We study the problem of sampling from a distribution under local differential privacy (LDP). Given a private distribution $P \in \mathcal{P}$, the goal is to generate a single sample from a distribution that remains close to P in f -divergence while satisfying the constraints of LDP. This task captures the fundamental challenge of producing realistic-looking data under strong privacy guarantees. While prior work by Park et al. (NeurIPS’24) focuses on global minimax-optimality across a class of distributions, we take a local perspective. Specifically, we examine the minimax risk in a neighborhood around a fixed distribution P_0 , and characterize its exact value, which depends on both P_0 and the privacy level. Our main result shows that the local minimax risk is determined by the global minimax risk when the distribution class $\mathcal{P}$ is restricted to a neighborhood around P_0 . To establish this, we (1) extend previous work from pure LDP to the more general functional LDP framework, and (2) prove that the globally optimal functional LDP sampler yields the optimal local sampler when constrained to distributions near P_0 . Building on this, we also derive a simple closed-form expression for the locally minimax-optimal samplers which does not depend on the choice of f -divergence. We further argue that this local framework naturally models private sampling with public data, where the public data distribution is represented by P_0 . In this setting, we empirically compare our locally optimal sampler to existing global methods, and demonstrate that it consistently outperforms global minimax samplers.

1 Introduction

Differential privacy (DP) [1] is the de facto standard for providing formal privacy guarantees in machine learning. A widely used variant, local differential privacy (LDP) [2], enables individuals to randomize their data on their own devices before sharing it with an untrusted curator. LDP ensures that the output of a randomized algorithm is statistically indistinguishable under arbitrary changes to the input, thereby limiting the privacy leakage of an individual’s data. By eliminating the need for a trusted curator, LDP has become particularly appealing to users and companies alike, and has seen wide deployment in practice—including systems developed by Google, Apple, and Microsoft[3–6].

Despite this progress, much of the LDP literature assumes that each user holds only a single data point. This assumption is often unrealistic, as modern devices routinely collect large volumes of data—such as images, messages, or time-series records [7]. To address this, a line of work known as

user-level DP assumes that each user holds a dataset of fixed size, with samples drawn i.i.d. from an underlying distribution [8–16]. However, the requirement of equal dataset size across users limits its practicality.

To overcome this limitation, *locally private sampling* was introduced by Husain et al. [7] and later extended by Park et al. [17]. This framework considers a more general setting where clients possess large datasets of varying sizes and aim to privately release another dataset that closely resembles their original data. Each client’s local dataset is modeled as an empirical distribution P , from which a single private sample is to be generated. Privacy is enforced by projecting P onto a set of distributions—called a *mollifier*—within a certain radius. While Husain et al. [7] construct the mollifier in an ad-hoc manner using a reference distribution, Park et al. [17] develop a principled minimax framework to identify the *optimal* mollifier, which corresponds to the worst-case data distribution (namely, degenerate distributions in the discrete case).

In this work, we extend the framework of Park et al. [17] to a public-private setting, where each individual also has access to a public dataset, represented by a distribution P_0 . From a theoretical perspective, our goal is to move beyond the pessimistic worst-case formulation and instead characterize optimal samplers through a local minimax lens. From a practical standpoint, this framework is motivated by the increasing relevance of personalized collaborative learning [18–20], in which individual models are trained collaboratively, and the public data may reflect (privatized) information shared by other users. Since real-world data distributions are rarely worst-case, our local formulation better captures achievable performance in practice.

Our local minimax framework aims to identify locally private samplers that minimize the f -divergence between the true distribution P and the sampling distribution, for all P within a neighborhood $N_\gamma(P_0)$ around P_0 , parameterized by a constant $\gamma \geq 1$. This neighborhood is defined via one of the strictest notions of distance between distributions: we say $Q \in N_\gamma(P_0)$ if $Q(x) \leq \gamma P_0(x)$ and $P_0(x) \leq \gamma Q(x)$ for all x . We will discuss this notion of neighborhood in details in Section 4.

Interestingly, we show that the local minimax-optimal private sampler under this notion of neighborhood is fully characterized by the global minimax-optimal sampler for a restricted class of data distributions. This relationship is most transparent in the context of functional LDP [21]. We therefore first extend the global minimax framework of Park et al. [17] from pure LDP to functional LDP, and then use it to derive the local minimax-optimal samplers.

Contributions. Our main contributions are as follows:

- We characterize the exact global minimax risk of private sampling under the functional LDP framework and derive optimal samplers for both continuous and discrete domains. Leveraging properties of functional LDP [21], our results generalize the pure-LDP framework of Park et al. [17] to approximate and Gaussian LDP (GLDP) notions (Section 3). As in their case, our optimal samplers are independent of the choice of f -divergence.
- We introduce a local minimax formulation to identify samplers that minimize f -divergence over all $P \in N_\gamma(P_0)$ for a given P_0 and $\gamma \geq 1$. We then show that this minimax risk is fully determined by the global minimax risk when the distribution class is restricted to a neighborhood around P_0 (Section 4). Building on our global minimax results, we derive closed-form expressions for the local minimax-optimal functional samplers. We further specialize this to express a pointwise optimal private sampler under pure LDP that achieves the local minimax risk (Section 5), enabling direct comparison with the global samplers of Park et al. [17].
- We numerically validate our local minimax samplers, showing they consistently—and often substantially—outperform the global samplers of Park et al. [17] across privacy regimes (Section 6). Figure 1 illustrates that our samplers yield distributions significantly closer to the original under both pure LDP and GLDP.¹

Related work. Locally private sampling was initially introduced by Husain et al. [7] and later formalized through a global minimax framework by Park et al. [17]. In their work, two fundamentally different families of minimax-optimal samplers were identified: one derived from the randomized

¹Experiment code is publicly available at https://github.com/hradghoukasian/private_sampling, with reproduction instructions provided in Appendix F. For completeness, the code is also included in the supplementary material.

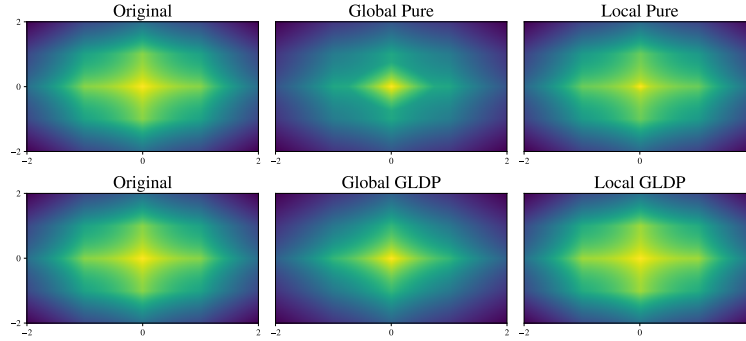


Figure 1: Comparison of global and local minimax-optimal sampler under pure LDP ($\varepsilon = 1$) and ν -GLDP ($\nu = 1.5$). Full details of this experiment are provided in Appendix E.1.

response mechanism, referred to as the *linear sampler*, and another based on clipping the data distribution, known as the *non-linear sampler*. They rigorously showed that the non-linear sampler is pointwise better than the linear one. In a similar vein, we develop two families of locally minimax-optimal pure LDP samplers and demonstrate that one is pointwise better than the other.

Sampling with public data was recently explored by Zamanlooy et al. [22], who study minimax-optimal samplers in the presence of a public prior. While related, their work differs in key ways: they focus on a global minimax formulation restricted to discrete domains and linear samplers (i.e., perturbing the data distribution linearly), whereas we study a local minimax problem, support both discrete and continuous settings, and consider arbitrary samplers. Moreover, the role of public data differs between the two approaches. Zamanlooy et al. [22] impose it as a hard constraint: their samplers are required to preserve the public prior exactly, ensuring that the public data remains unchanged while privacy is enforced only on the private data. By contrast, we use the public distribution to define a local neighborhood, which in turn yields a local minimax formulation applicable to general sample spaces. A key implication of this distinction is that in their setting, access to public data does not necessarily reduce the minimax risk, whereas in ours it does.

Private sampling under central DP was introduced by Raskhodnikova et al. [23], who proposed algorithms for generating private samples from k -ary and product Bernoulli distributions. They showed that sampling can incur significantly lower privacy costs than private learning for some range of parameters. This was extended to broader distribution families [24–26] and to multi-sample settings [27].

Private sampling has also gained traction in fine-tuning large language models [28, 29], and is closely related to private generative modeling. However, most prior work in this area focuses on central DP [30–37], assumes single-record LDP [38–41], or targets specific estimation tasks [42].

Finally, we remark that local minimax formulations have also been explored in the context of estimation and distribution learning under DP [43–48]. In particular, Chen and Özgür [43] analyze discrete distribution estimation under LDP by considering neighborhoods around a reference distribution. Similarly, McMillan et al. [44] and Duchi and Ruan [45] develop locally minimax estimators for one-dimensional parameters in the central and local models, respectively, while Asi and Duchi [46] extend this framework to function estimation. More recently, Feldman et al. [47] utilize the notion of neighborhoods around a distribution in the context of nonparametric density estimation. Although these works focus primarily on estimating statistical functionals or learning distributions, they are conceptually distinct from our setting, which centers on private sampling. Nonetheless, our notion of neighborhood in the local minimax formulation is closely aligned with that of Feldman et al. [47].

Notation. Let $\mathcal{P}(\mathcal{X})$ denote the set of all probability distributions over the sample space $\mathcal{X}$, and let $\mathcal{C}(\mathbb{R}^n)$ denote the set of all continuous probability distributions on $\mathbb{R}^n$. For $P \in \mathcal{C}(\mathbb{R}^n)$, we denote its probability density function (PDF) by the lowercase letter p . For any integer $k \in \mathbb{N}$, we define $[k] := \{1, 2, \dots, k\}$. The notation $\text{clip}(x; s_1, s_2) := \max\{s_1, \min\{s_2, x\}\}$ refers to the clipping function that bounds x between s_1 and s_2 . The convex conjugate of a function $g : D \subseteq \mathbb{R} \rightarrow \mathbb{R}$ is the function $g^* : \mathbb{R} \rightarrow (-\infty, +\infty]$ defined by $g^*(y) := \sup_{x \in D} \{xy - g(x)\}$ for $y \in \mathbb{R}$. We denote the n -dimensional Laplace distribution with mean $m \in \mathbb{R}^n$ and scale $b > 0$ by $\mathcal{L}(m, b)$, and

its density at $x \in \mathbb{R}^n$ by $\mathcal{L}(x \mid m, b)$. For the Gaussian case, $\mathcal{N}(m, \Sigma)$ denotes the n -dimensional normal distribution with mean $m \in \mathbb{R}^n$ and covariance matrix Σ .

2 Preliminaries

In this section, we formally introduce definitions and concepts necessary for the subsequent sections.

Suppose a client holds a distribution $P \in \tilde{\mathcal{P}} \subset \mathcal{P}(\mathcal{X})$ over a sample space $\mathcal{X}$, and wishes to release a sample that resembles being drawn from P while preserving privacy. To ensure privacy, the client applies a private sampler $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$, which maps the input distribution P to a privatized distribution $\mathbf{Q}(P) \in \mathcal{P}(\mathcal{X})$. The released sample is then drawn from $\mathbf{Q}(P)$, thereby preserving privacy while approximating the original distribution. We equivalently view $\mathbf{Q}(P)$ as the conditional distribution $\mathbf{Q}(\cdot \mid P)$.

Definition 2.1 (Approximate LDP). Let $\varepsilon \geq 0$ and $\delta \in [0, 1]$. A sampler $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$ is said to satisfy (ε, δ) -local differential privacy $((\varepsilon, \delta)$ -LDP) if, for every pair of input distributions $P, P' \in \tilde{\mathcal{P}}$ and every measurable set $A \subseteq \mathcal{X}$, we have

$$\mathbf{Q}(A \mid P) \leq e^\varepsilon \mathbf{Q}(A \mid P') + \delta.$$

When $\delta = 0$, $\mathbf{Q}$ is ε -LDP sampler [7], also referred to as *pure LDP*. We let $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}$ and $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}$ denote the sets of all ε -LDP and (ε, δ) -LDP samplers, respectively.

Definition 2.1 ensures privacy by requiring that the output of the sampler remains indistinguishable under any arbitrary changes in the input. In contrast, functional LDP [21] interprets privacy through the lens of hypothesis testing: given two inputs, let P and Q be the corresponding output distributions. The privacy guarantee is quantified by how hard it is for an adversary to distinguish P from Q , which can be framed as a binary hypothesis testing problem [49, 50]:

$$H_0 : \text{output} \sim P \quad \text{vs} \quad H_1 : \text{output} \sim Q.$$

Given a rejection rule $\phi : \mathcal{X} \rightarrow [0, 1]$, the Type I and Type II error rates are defined as

$$a_\phi := \mathbb{E}_P[\phi], \quad b_\phi := 1 - \mathbb{E}_Q[\phi],$$

respectively.

Definition 2.2 (Trade-off function). Let P and Q be probability measures in $\mathcal{P}(\mathcal{X})$. The *trade-off function* $T(P, Q) : [0, 1] \rightarrow [0, 1]$ is defined as

$$T(P, Q)(u) = \inf_{\phi} \{ b_\phi : a_\phi \leq u \},$$

where the infimum is taken over all measurable rejection rules ϕ .

Notice that $T(P, Q)(u)$ represents the smallest achievable Type II error subject to Type I error being at most u . A function g is called a trade-off function if $g = T(P, Q)$ for some distributions P and Q .

Definition 2.3 (Functional LDP [21, 51]). Given a trade-off function $g : [0, 1] \rightarrow [0, 1]$, a sampler $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$ is said to satisfy *g-functional local differential privacy* (g -FLDP) if for every pair of input distributions $P, P' \in \tilde{\mathcal{P}}$ and every $u \in [0, 1]$,

$$T(\mathbf{Q}(\cdot \mid P), \mathbf{Q}(\cdot \mid P'))(u) \geq g(u).$$

We let $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}$ denote the sets of all samplers $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$ satisfying g -FLDP.

Note that functional LDP generalizes both approximate and pure LDP. Specifically, for any privacy parameters (ε, δ) , there exists a trade-off function $g_{\varepsilon, \delta}$ such that $g_{\varepsilon, \delta}$ -FLDP is equivalent to (ε, δ) -LDP [21, Proposition 3]. This function is given by $g_{\varepsilon, \delta}(\theta) = \max \{0, 1 - \delta - e^\varepsilon \theta, e^{-\varepsilon}(1 - \delta - \theta)\}$. The case of pure LDP corresponds to the special case $g_\varepsilon := g_{\varepsilon, 0}$. Another notable instance of g -FLDP is ν -GLDP, where the trade-off function is given by $G_\nu(\theta) = \Phi(\Phi^{-1}(1 - \theta) - \nu)$, with Φ denoting the standard normal CDF [21, 51].

In order to quantify utility loss incurred by sampler $\mathbf{Q}$, we use the f -divergence between the original distribution P and the resulting sampling distribution $\mathbf{Q}(P)$, defined as follows.

Definition 2.4 (f -divergence [52, 53]). Let $f : (0, \infty) \rightarrow \mathbb{R}$ be a convex function with $f(1) = 0$. The f -divergence between $P, Q \in \mathcal{P}(\mathcal{X})$ with $P \ll Q$ is defined as $D_f(P \| Q) := \mathbb{E}_Q \left[f \left(\frac{dP}{dQ} \right) \right]$.

The above definition can be extended to cases where $P \ll Q$ is no longer satisfied as in Appendix B.1. Common examples of f -divergences include KL divergence, total variation distance, and squared Hellinger distance. We denote total variation distance between $P, Q \in \mathcal{P}(\mathcal{X})$ by $D_{\text{TV}}(P, Q)$. A key instance in this work is the E_γ -divergence (or hockey-stick divergence), defined as the f -divergence with $f(t) = \max\{t - \gamma, 0\}$ for $\gamma \geq 1$ [54]. This divergence is particularly relevant as it defines the local neighborhood in our minimax formulation. We denote E_γ -divergence of two distributions P and Q as $E_\gamma(P \| Q)$.

Our most general result, covering both continuous and discrete spaces, requires a technical assumption on distributions, stated below and discussed further in Section 3.

Definition 2.5 (Decomposability [17]). Let $\alpha \in (0, 1)$ and $t, u \in \mathbb{N}$ with $t > u$. A probability measure $\mu \in \mathcal{P}(\mathcal{X})$ is (α, t, u) -decomposable if there exist sets $A_1, \dots, A_t \subseteq \mathcal{X}$ such that $\mu(A_i) = \alpha$ for all $i \in [t]$, and for every $x \in \mathcal{X}$, the number of sets A_i containing x is at most u .

Establishing the exact optimal minimax risk in our proofs involves deriving matching upper bounds (the achievability part) and lower bounds (the converse part). For a unified proof over general sample space $\mathcal{X}$, the converse part relies on $(\alpha, t, 1)$ -decomposability; namely, $\mathcal{X}$ contains t disjoint measurable subsets, each of measure α . We note that this mild condition holds naturally for absolutely continuous distributions and for the uniform reference measure in the finite case.

3 Global minimax-optimal samplers under functional LDP

In this section, we study the global minimax samplers under the general framework of functional LDP and derive compact characterizations of the optimal samplers. These global samplers will serve as proxies for constructing local minimax-optimal samplers in subsequent sections.

Using f -divergence as the utility measure, we define the global minimax risk as

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f) := \inf_{\mathcal{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathcal{Q}(P)). \quad (1)$$

As noted by Park et al. [17], this risk becomes vacuous in the continuous setting if the distribution class $\tilde{\mathcal{P}}$ is unrestricted. To address this, they consider a restricted class of distributions defined as

$$\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, h} := \{P \in \mathcal{C}(\mathbb{R}^n) : c_1 h(x) \leq p(x) \leq c_2 h(x), \quad \forall x \in \mathbb{R}^n\}, \quad (2)$$

for a reference function $h : \mathcal{X} \rightarrow [0, \infty)$ satisfying $\int_{\mathbb{R}^n} h(x) dx < \infty$, and constants $c_2 > c_1 \geq 0$. This class captures a broad range of practical distributions, including mixtures of Gaussians and Laplace distributions. See the example below and further discussion in Appendices E.1 and E.2.

Example 3.1. For $n, k \in \mathbb{N}$, define $\tilde{\mathcal{P}}_{\mathcal{L}} = \left\{ \sum_{i=1}^k \lambda_i \mathcal{L}(m_i, b) : \lambda_i \geq 0, \sum_{i=1}^k \lambda_i = 1, \|m_i\|_1 \leq 1 \right\}$ and $\tilde{\mathcal{P}}_{\mathcal{N}} = \left\{ \sum_{i=1}^k \lambda_i \mathcal{N}(m_i, \sigma^2 I_n) : \lambda_i \geq 0, \sum_{i=1}^k \lambda_i = 1, \|m_i\|_2 \leq 1 \right\}$. Here $m_i, x \in \mathbb{R}^n$ and I_n denotes the $n \times n$ identity matrix. It can be verified that $\tilde{\mathcal{P}}_{\mathcal{N}} \subseteq \tilde{\mathcal{P}}_{0, 1, h_{\mathcal{N}}}$ and $\tilde{\mathcal{P}}_{\mathcal{L}} \subseteq \tilde{\mathcal{P}}_{e^{-1/b}, e^{1/b}, h_{\mathcal{L}}}$, where $h_{\mathcal{N}}(x) = \frac{1}{(2\pi\sigma^2)^{\frac{n}{2}}} \exp\left(-\frac{[\max(0, \|x\|_2 - 1)]^2}{2\sigma^2}\right)$ and $h_{\mathcal{L}}(x) = \mathcal{L}(x | 0, b)$.

Without loss of generality, we assume that $\int_{\mathbb{R}^n} h(x) dx = 1$ and $0 \leq c_1 < 1 < c_2$. Moreover, for technical reasons—explained in Appendix D.2—we also assume that $\frac{c_2 - c_1}{1 - c_1} \in \mathbb{N}$. This assumption is not restrictive and holds in many practical settings, e.g., Gaussian mixtures (see Example 3.1). When it is not satisfied, one can slightly increase c_2 or decrease c_1 to enforce the condition, resulting in a superset of the original distribution class, i.e., $\tilde{\mathcal{P}} \subseteq \tilde{\mathcal{P}}_{c_1, c_2, h}$.

While many of our examples focus on continuous domains, our analysis applies more generally to any sample space $\mathcal{X}$, provided the distributions are absolutely continuous with respect to a reference measure μ . In this case, we define $\tilde{\mathcal{P}}_{c_1, c_2, \mu} := \{P \in \mathcal{P}(\mathcal{X}) : P \ll \mu, c_1 \leq \frac{dP}{d\mu} \leq c_2, \mu\text{-a.e.}\}$, which we take as the generalized distribution class. Thus, we make the following assumption.

Assumption 3.2. We assume $\int_{\mathbb{R}^n} h(x) dx = 1$ and $\mu(\mathcal{X}) = 1$ for the classes $\tilde{\mathcal{P}}_{c_1, c_2, h}$ and $\tilde{\mathcal{P}}_{c_1, c_2, \mu}$, respectively. Additionally, we assume that $0 \leq c_1 < 1 < c_2$, and $\frac{c_2 - c_1}{1 - c_1} \in \mathbb{N}$.

The following proposition shows that an additional condition is required to exclude trivial samplers.

Proposition 3.3. Let g be a trade-off function. Under Assumption 3.2, if $1 + g^*(-e^\beta) > \frac{(c_2 - c_1 e^\beta)(1 - c_1)}{c_2 - c_1}$ for all $\beta \geq 0$, then $\mathbf{Q}(P) = P$ satisfies g -FLDP.

In view of this result, we assume throughout that for any given trade-off function g the following condition holds:

$$\exists \beta \geq 0 \text{ such that } 1 + g^*(-e^\beta) \leq \frac{(c_2 - c_1 e^\beta)(1 - c_1)}{c_2 - c_1}. \quad (3)$$

Now, we are in order to state our first technical result.

Theorem 3.4. Let $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, h}$. Under Assumption 3.2, the sampler $\mathbf{Q}_{c_1, c_2, h, g}^*$ defined as a continuous distribution whose density is given by

$$q_g^*(P)(x) := \lambda_{c_1, c_2, g}^* p(x) + (1 - \lambda_{c_1, c_2, g}^*) h(x), \quad \lambda_{c_1, c_2, g}^* = \inf_{\beta \geq 0} \frac{e^\beta + \frac{c_2 - c_1}{1 - c_1} (1 + g^*(-e^\beta)) - 1}{(1 - c_1) e^\beta + c_2 - 1}, \quad (4)$$

satisfies g -FLDP, and is minimax-optimal under any f -divergence, that is

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, h, g}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1), \quad (5)$$

for $r_1 = \frac{c_1}{1 - (1 - c_1) \lambda_{c_1, c_2, g}^*}$ and $r_2 = \frac{c_2}{(c_2 - 1) \lambda_{c_1, c_2, g}^* + 1}$.

We note that the non-triviality condition in Proposition 3.3 is equivalent to $\lambda_{c_1, c_2, g}^* \leq 1$, which allows a natural interpretation of the sampler $\mathbf{Q}_{c_1, c_2, h, g}^*$: It samples from P with probability $\lambda_{c_1, c_2, g}^*$ and from the reference distribution with probability $1 - \lambda_{c_1, c_2, g}^*$. This probabilistic mixture reflects the privacy–utility trade-off: a larger $\lambda_{c_1, c_2, g}^*$ yields outputs closer to P , while sampling from h introduces the randomness required for privacy. The same mixture structure yields a minimax-optimal sampler for the discrete case where $\mathcal{X} = [k]$ and $\tilde{\mathcal{P}} = \mathcal{P}([k])$.

Theorem 3.5. Let $\tilde{\mathcal{P}} = \mathcal{P}([k])$ for some $k \in \mathbb{N}$, and let μ_k denote the uniform distribution on $[k]$. Then the sampler defined as

$$\mathbf{Q}_{k, g}^*(P) := \lambda_{k, g}^* P + (1 - \lambda_{k, g}^*) \mu_k, \quad \lambda_{k, g}^* = \inf_{\beta \geq 0} \frac{e^\beta + k(1 + g^*(-e^\beta)) - 1}{e^\beta + k - 1},$$

satisfies g -FLDP, and is minimax-optimal under any f -divergence, that is

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{k, g}^*(P)) = \mathcal{R}(\mathcal{Q}_{[k], \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f).$$

Theorems 3.4 and 3.5 provide optimal samplers for continuous and finite sample spaces, respectively. The next theorem addresses the same task for general sample spaces under the assumption of decomposability.

Theorem 3.6. Let $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$ under Assumption 3.2, and suppose μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1 - c_1}{c_2 - c_1}$. Then the sampler defined as $\mathbf{Q}_{c_1, c_2, \mu, g}^*(P) = \lambda_{c_1, c_2, g}^* P + (1 - \lambda_{c_1, c_2, g}^*) \mu$ satisfies g -FLDP, and is minimax-optimal with respect to any f -divergence, where $\lambda_{c_1, c_2, g}^*$ is defined in Theorem 3.4.

Similar to Theorem 3.6, all subsequent results can be extended to a general sample space. For clarity of presentation, however, we state them in the continuous setting, while all proofs are given in full generality.

Having established the optimal sampler for general g -FLDP, we present specific instantiations of the result for two widely studied settings: pure LDP and ν -GLDP. First, we derive the optimal ε -LDP sampling algorithm by setting $g = g_\varepsilon$ in Theorem 3.4, as stated in the following corollary.

Corollary 3.7. Let $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, h}$. Under Assumption 3.2, the sampler $\mathbf{Q}_{c_1, c_2, h, g_\varepsilon}^*$ defined as a continuous distribution whose density is given by $q_{g_\varepsilon}^*(P)(x) := \lambda_{c_1, c_2, g_\varepsilon}^* p(x) + (1 - \lambda_{c_1, c_2, g_\varepsilon}^*) h(x)$,

with $\lambda_{c_1, c_2, g_\varepsilon}^* = \frac{e^\varepsilon - 1}{(1 - c_1)e^\varepsilon + c_2 - 1}$, satisfies g_ε -FLDP, and is minimax-optimal with respect to any f -divergence, that is

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, h, g_\varepsilon}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, g_\varepsilon}, \tilde{\mathcal{P}}, f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1), \quad (6)$$

for $r_1 = c_1 \cdot \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{c_2 - c_1}$, and $r_2 = \frac{c_2}{c_2 - c_1} \cdot \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{e^\varepsilon}$.

Indeed, by framing ε -LDP as a special case of functional LDP, Corollary 3.7 reproduces the minimax risk of the optimal ε -LDP sampler in [17, Theorem 3.3]. Next, we focus on the ν -GLDP as a special case of functional LDP.

Corollary 3.8 (ν -GLDP as a special case). *Let $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, h}$. Under Assumption 3.2, the sampler (4) for $g = G_\nu$ belongs to $\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, G_\nu}$ and is minimax-optimal with respect to any f -divergence, that is,*

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{G_\nu}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, G_\nu}, \tilde{\mathcal{P}}, f). \quad (7)$$

The explicit description of the optimal sampler under ν -GLDP is given in Appendix D.6.

As outlined in the introduction, the global minimax formulation under functional LDP serves as the basis for our local analysis. We now turn to the local minimax setting, first under general g -FLDP (Section 4) and then under pure LDP (Section 5), leveraging Theorem 3.4 and Corollary 3.7, respectively.

4 Local minimax-optimal sampling under functional LDP

In this section, we begin by formalizing the local minimax setting using a neighborhood structure induced by the E_γ -divergence. We then develop a general framework for identifying local minimax-optimal samplers over arbitrary sample spaces.

The local minimax-optimal sampling problem adopts the same structural form as the global minimax problem, with a key distinction: instead of selecting the worst-case distribution from the entire space of the general class $\tilde{\mathcal{P}}$, the search is restricted to a neighborhood around a fixed distribution P_0 . This reference distribution P_0 may represent, for example, a publicly available dataset held by an individual or shared collaboratively in a distributed setting. Formally, the local minimax problem under g -FLDP is defined as:

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, N_\gamma(P_0), f) := \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}} \sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}(P)), \quad (8)$$

where $N_\gamma(P_0)$ denotes a neighborhood around P_0 , recently adopted by Feldman et al. [47], and is defined for any $\gamma \geq 1$ as

$$N_\gamma(P_0) := \{P \in \mathcal{P}(\mathcal{X}) : E_\gamma(P \| P_0) = E_\gamma(P_0 \| P) = 0\}. \quad (9)$$

The zero- E_γ neighborhood generalizes the classical notion of total variation distance. Specifically, for $\gamma \geq 1$, $E_\gamma(P \| Q) = 0$ implies $D_{\text{TV}}(P, Q) \leq 1 - \frac{1}{\gamma}$ [55, Proposition 4]. Moreover, we can straightforwardly extend this neighborhood to the more general case by replacing $E_\gamma(P \| P_0) = 0$ with $E_\gamma(P \| P_0) \leq \zeta$, meaning the likelihood ratio lies in $[1/\gamma, \gamma]$ with probability at least $1 - \zeta$ under P_0 . This, in turn, enables further generalization to any f -divergence with a twice differentiable function f , which is a common trick in information theory [56–59]. As in the global case (Section 3), we assume $\gamma \in \mathbb{N}$ for the same technical reasons. Theorem 4.1 presents the corresponding locally minimax-optimal sampler, connecting the global minimax analysis of Section 3 with the local setting under functional LDP.

Theorem 4.1. *Let P_0 be a continuous distribution on $\mathbb{R}^n$. Define $N_\gamma(P_0)$ as in (9), and assume $N_\gamma(P_0) \subseteq \tilde{\mathcal{P}}$. Let $\mathbf{Q}_g^*$ denote the global sampler from (4), instantiated with parameters $c_1 = \frac{1}{\gamma}$, $c_2 = \gamma$, and $h = p_0$. We define the sampler $\mathbf{Q}_{g, N_\gamma(P_0)}^*$ as*

$$\mathbf{Q}_{g, N_\gamma(P_0)}^*(P) := \begin{cases} \mathbf{Q}_g^*(P), & \text{if } P \in N_\gamma(P_0), \\ \mathbf{Q}_g^*(\hat{P}), & \text{otherwise,} \end{cases}$$

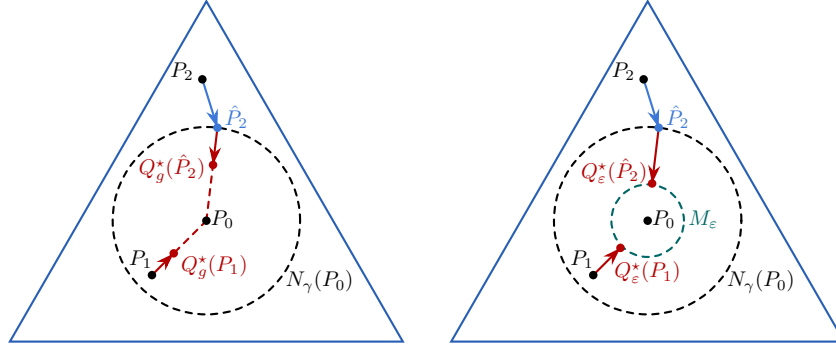


Figure 2: Illustrations of the optimal samplers described in Theorems 4.1 and 5.1. **Left:** $\mathbf{Q}_{g, N_\gamma(P_0)}^*$ for functional LDP. **Right:** $\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*$ for pure LDP. Following [17, Proposition 3.4], M_ε is defined as $M_\varepsilon := \{Q \in \mathcal{C}(\mathbb{R}^n) : \frac{\gamma+1}{\gamma+e^\varepsilon} p_0(x) \leq q(x) \leq \frac{(\gamma+1)e^\varepsilon}{\gamma+e^\varepsilon} p_0(x), \forall x \in \mathbb{R}^n\}$.

where $\hat{P} \in N_\gamma(P_0)$ is a distribution that minimizes $D_f(P \| P')$ over all $P' \in N_\gamma(P_0)$. Then $\mathbf{Q}_{g, N_\gamma(P_0)}^* \in \mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, g}$ and is locally minimax-optimal under any f -divergence, that is

$$\sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_{g, N_\gamma(P_0)}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, g}, N_\gamma(P_0), f).$$

Theorem 4.1 establishes that the local minimax-optimal sampler for the problem in (8) coincides with the global solution in (1) when the universe is restricted to the N_γ neighborhood around P_0 . In this sense, the global solution presented in Theorem 3.4 directly yields the local solution in Theorem 4.1.

By instantiating Theorem 4.1 with $g = g_\varepsilon$, we obtain an ε -LDP sampler that is locally minimax-optimal. However, as we demonstrate in the next section, this sampler can still be improved in a *pointwise* sense.

5 Local minimax-optimal sampling under pure LDP

In this section, we turn our attention to pure-LDP sampling and develop a new framework that is fundamentally different from the approach presented in the previous section. Our goal is to identify a sampler that is *pointwise* better than the one obtained from Theorem 4.1 with $g = g_\varepsilon$. The construction of such sampler is presented in the next theorem.

Theorem 5.1. Let P_0 be a continuous distribution on $\mathbb{R}^n$ with density p_0 and $N_\gamma(P_0)$ be the neighborhood around P_0 defined in (9). Assuming $N_\gamma(P_0) \subseteq \tilde{\mathcal{P}}$, let $\mathbf{Q}_\varepsilon^* \in \mathcal{Q}_{\mathbb{R}^n, N_\gamma(P_0), \varepsilon}$ be a sampler with $\mathbf{Q}_\varepsilon^*(P)$ having density

$$q(x) = \text{clip} \left(\frac{1}{r_P} p(x); \frac{\gamma+1}{\gamma+e^\varepsilon} p_0(x), \frac{(\gamma+1)e^\varepsilon}{\gamma+e^\varepsilon} p_0(x) \right),$$

where r_P being the normalizing constant ensuring $\int_{\mathbb{R}^n} q(x) dx = 1$. Furthermore, define the extended sampler $\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*$ by

$$\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*(P) := \begin{cases} \mathbf{Q}_\varepsilon^*(P), & \text{if } P \in N_\gamma(P_0), \\ \mathbf{Q}_\varepsilon^*(\hat{P}), & \text{otherwise,} \end{cases}$$

where $\hat{P} \in N_\gamma(P_0)$ is a distribution that minimizes $D_f(P \| P')$ over all $P' \in N_\gamma(P_0)$. Then, we have $\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^* \in \mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, \varepsilon}$ and $\sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, \varepsilon}, N_\gamma(P_0), f)$.

Similar to Theorem 4.1, the above theorem defines a sampler that distinguishes between the cases $P \in N_\gamma(P_0)$ and $P \notin N_\gamma(P_0)$. In the latter case, both theorems are aligned in that the sampling distribution depends on $\hat{P}$, the closest distribution to P within $N_\gamma(P_0)$. However, for the case

$P \in N_\gamma(P_0)$, the two constructions differ fundamentally. Theorem 4.1 assigns a *linear* sampler, whereas Theorem 5.1 carefully constructs a non-linear sampler tailored to this regime. In fact, this non-linear sampler is obtained by projecting P onto a mollifier M_ε ; see proof of the theorem in Appendix D.8 and Figure 2 for a visual representation. This projection is carried out through a clipping operation, which guarantees that the resulting distribution $\mathbf{Q}_\varepsilon^*(P)$ satisfies the LDP constraint while maintaining a close match to the original distribution P in terms of f -divergence. This approach yields a nonlinear transformation of P that carefully balances privacy constraints and utility preservation. It can be verified—using a mutatis mutandis adaptation of [17, Proposition C.5]—that the non-linear sampler pointwise outperforms the linear one under the same privacy guarantees; that is, for any $\varepsilon \geq 0$ and $P \in N_\gamma(P_0)$

$$D(P \parallel \mathbf{Q}_\varepsilon^*(P)) \leq D(P \parallel \mathbf{Q}_{g_\varepsilon}^*(P)).$$

In fact, the non-linear sampler outperforms the linear one because it is not only worst-case optimal but also instance-optimal: for each input distribution P in our distribution class and any f -divergence D_f , it minimizes $D_f(P \parallel \mathbf{Q})$ over all admissible $\mathbf{Q}$ satisfying ε -LDP (see Proposition D.2). In contrast, the linear sampler uses a fixed transformation for all input distributions. This instance-specific optimization explains the empirical advantage of the non-linear sampler. As a result, the overall sampler in Theorem 5.1 is pointwise better than the one in Theorem 4.1. This intuition is formalized in the result below.

Proposition 5.2. *Let $\tilde{\mathcal{P}}$ be the global universe and P_0 a probability measure on $\mathbb{R}^n$, with $N_\gamma(P_0) \subseteq \tilde{\mathcal{P}}$ for $N_\gamma(P_0)$ defined as in (9). Let $\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*$ be the optimal ε -LDP sampler from Theorem 5.1, and $\mathbf{Q}_{g_\varepsilon, N_\gamma(P_0)}^*$ the instantiation of Theorem 4.1 with $g = g_\varepsilon$. Then, for all $P \in N_\gamma(P_0)$,*

$$D_f(P \parallel \mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*) \leq D_f(P \parallel \mathbf{Q}_{g_\varepsilon, N_\gamma(P_0)}^*).$$

6 Numerical results

In this section, we numerically compare the worst-case f -divergence of our locally minimax sampler against the globally optimal sampler of [17] under the ε -LDP setting. Experiments span both finite and continuous domains, evaluating KL divergence, total variation distance, and squared Hellinger distance across $\varepsilon \in \{0.1, 0.5, 1, 2\}$. Additional results under ν -GLDP appear in Appendix E.5 and E.6. In addition, we report complementary results for the continuous case under pure LDP in Appendix E.7.

6.1 Finite sample space

We compare local and global minimax-optimal samplers in the finite setting $\mathcal{X} = [k]$, where the global class is $\tilde{\mathcal{P}}_{\text{global}} = \mathcal{P}([k]) = \tilde{\mathcal{P}}_{0,k,\mu_k}$ for the uniform measure μ_k . The local neighborhood is defined as $\tilde{\mathcal{P}}_{\text{local}} = \mathcal{N}_\gamma(\mu_k) = \tilde{\mathcal{P}}_{\frac{1}{\gamma},\gamma,\mu_k}$ with $\gamma = \frac{k}{2} - 1$, ensuring that all local neighborhood assumptions are satisfied and $\mathcal{N}_\gamma(\mu_k) \subseteq \mathcal{P}([k])$. Indeed, both global and local worst-case f -divergences can be computed in closed form: the global risk is given by [17, Theorem 3.1], while the local risk is derived from a finite-space version of Theorem 5.1 (see Appendix E.3). Figure 3 compares the two for $k = 20$; additional results for other k values are provided in Appendix E.3. The local minimax-optimal sampler consistently outperforms the global minimax-optimal sampler across all ε and f -divergences.

6.2 Continuous sample space

In the continuous setting with $\mathcal{X} = \mathbb{R}$, we fix the universe $\tilde{\mathcal{P}}_{\text{local}}$ and evaluate the empirical worst-case f -divergence over 100 randomly generated client distributions $P_1, \dots, P_{100} \in \tilde{\mathcal{P}}_{\text{local}}$. Each P_j represents a client and is constructed as a mixture of a random number of one-dimensional Laplace components with scale parameter $b = 1$; the complete procedure for generating these distributions is described in detail in Appendix E.4.

We define the local and global universes as $\tilde{\mathcal{P}}_{\text{local}} = \tilde{\mathcal{P}}_{1/3,3,h_{\mathcal{L}}}$ and $\tilde{\mathcal{P}}_{\text{global}} = \tilde{\mathcal{P}}_{1/9,9,h_{\mathcal{L}}}$, where $h_{\mathcal{L}}$ is the density of a Laplace distribution with mean zero and scale $b = 1$. Since $\tilde{\mathcal{P}}_{\text{local}} \subseteq \tilde{\mathcal{P}}_{\text{global}}$, every client distribution P_j also belongs to the global universe; thus the input distributions are

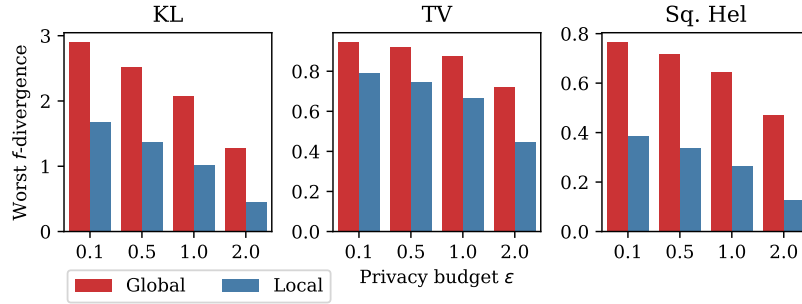


Figure 3: Theoretical worst-case f -divergences of global and local minimax samplers under the pure LDP setting with uniform reference distribution μ_k over finite space ($k = 20$).

identical for both samplers. We evaluate the empirical worst-case divergence of each sampler as $\max_{j \in [100]} D_f(P_j \| Q(P_j))$. The local minimax sampler follows Theorem 5.1, while the global sampler is the optimal sampler from [17, Theorem 3.3]. As illustrated in Figure 4, the local sampler consistently achieves lower worst-case f -divergence than its global counterpart across all f -divergences and privacy levels ϵ .

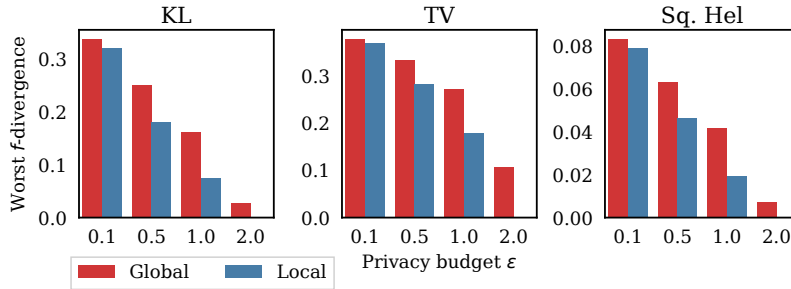


Figure 4: Empirical worst-case f -divergences of global and local minimax samplers under the pure LDP setting, over 100 experiments on a 1-D Laplace mixture.

7 Limitations and future direction

Our main contribution is the development of a minimax-optimal private sampler, validated through a series of experiments on synthetic data. However, the absence of evaluations on high-dimensional real-world datasets limits our understanding of its practical utility. A key reason for this is the computational complexity of our sampler (in particular the clipping function), which poses challenges for scalability in high-dimensional settings. Addressing this limitation—by developing more efficient implementations of our samplers or approximate variants, potentially leveraging techniques such as MCMC—is an important direction for future work.

Beyond computational considerations, another important direction for future work is to generalize the formulation of the local neighborhood. The neighborhood in our local minimax formulation is defined using the E_γ -divergence: $P \in N_\gamma(P_0)$ if $E_\gamma(P \| P_0) = E_\gamma(P_0 \| P) = 0$, as in [47]. A natural extension is $N_{\gamma,\zeta}(P_0) := \{P \in \mathcal{P}(\mathcal{X}) : E_\gamma(P \| P_0) \leq \zeta \text{ and } E_\gamma(P_0 \| P) \leq \zeta\}$. The use of E_γ -divergence in our formulation is particularly useful, as it provides a foundation for broader neighborhood definitions. In particular, it can be extended to neighborhoods based on general f -divergences with twice differentiable f , as $N_{f,\zeta}(P_0) := \{P \in \mathcal{P}(\mathcal{X}) : D_f(P \| P_0) \leq \zeta \text{ and } D_f(P_0 \| P) \leq \zeta\}$ [56–59]. This approach provides a principled strategy for characterizing the local minimax solution over the broad and flexible class of distributional neighborhoods $N_{f,\zeta}(P_0)$. We leave this generalization for future work. Furthermore, this work focuses on the setting where each client releases a single sample. Extending to the case of multiple samples per client is a natural direction for future work.

Finally, the broader impacts of our work is presented in Appendix G.

Acknowledgments

This work was supported in part by the Natural Sciences and Engineering Research Council of Canada (NSERC).

References

- [1] Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Calibrating noise to sensitivity in private data analysis. In *Theory of Cryptography: Third Theory of Cryptography Conference, TCC 2006, New York, NY, USA, March 4-7, 2006. Proceedings 3*, pages 265–284. Springer, 2006.
- [2] Shiva Prasad Kasiviswanathan, Homin K Lee, Kobbi Nissim, Sofya Raskhodnikova, and Adam Smith. What can we learn privately? *SIAM Journal on Computing*, 40(3):793–826, 2011.
- [3] Úlfar Erlingsson, Vasyl Pihur, and Aleksandra Korolova. Rappor: Randomized aggregatable privacy-preserving ordinal response. In *Proceedings of the 2014 ACM SIGSAC conference on computer and communications security*, pages 1054–1067, 2014.
- [4] Andrea Bittau, Úlfar Erlingsson, Petros Maniatis, Ilya Mironov, Ananth Raghunathan, David Lie, Mitch Rudominer, Ushasree Kode, Julien Tinnes, and Bernhard Seefeld. Prochlo: Strong privacy for analytics in the crowd. In *Proceedings of the 26th symposium on operating systems principles*, pages 441–459, 2017.
- [5] Differential privacy team Apple. Learning with privacy at scale, 2017.
- [6] Bolin Ding, Janardhan Kulkarni, and Sergey Yekhanin. Collecting telemetry data privately. *Advances in Neural Information Processing Systems*, 30, 2017.
- [7] Hisham Husain, Borja Balle, Zac Cranko, and Richard Nock. Local differential privacy for sampling. In *International Conference on Artificial Intelligence and Statistics*, pages 3404–3413. PMLR, 2020.
- [8] Alexander Kent, Thomas B Berrett, and Yi Yu. Rate optimality and phase transition for user-level local differential privacy. *arXiv preprint arXiv:2405.11923*, 2024.
- [9] Yulian Mao, Qingqing Ye, Haibo Hu, Qi Wang, and Kai Huang. Privshape: Extracting shapes in time series under user-level local differential privacy. In *2024 IEEE 40th International Conference on Data Engineering (ICDE)*, pages 1739–1751. IEEE, 2024.
- [10] Raef Bassily and Ziteng Sun. User-level private stochastic convex optimization with optimal rates. In *International Conference on Machine Learning*, pages 1838–1851. PMLR, 2023.
- [11] Jayadev Acharya, Yuhan Liu, and Ziteng Sun. Discrete distribution estimation under user-level local differential privacy. In *International Conference on Artificial Intelligence and Statistics*, pages 8561–8585. PMLR, 2023.
- [12] Ruiquan Huang, Huanyu Zhang, Luca Melis, Milan Shen, Meisam Hejazinia, and Jing Yang. Federated linear contextual bandits with user-level differential privacy. In *International Conference on Machine Learning*, pages 14060–14095. PMLR, 2023.
- [13] Antonious M Girgis, Deepesh Data, and Suhas Diggavi. Distributed user-level private mean estimation. In *2022 IEEE International Symposium on Information Theory (ISIT)*, pages 2196–2201. IEEE, 2022.
- [14] Daniel Levy, Ziteng Sun, Kareem Amin, Satyen Kale, Alex Kulesza, Mehryar Mohri, and Ananda Theertha Suresh. Learning with user-level privacy. *Advances in Neural Information Processing Systems*, 34:12466–12479, 2021.
- [15] Badih Ghazi, Ravi Kumar, and Pasin Manurangsi. User-level differentially private learning via correlated sampling. *Advances in Neural Information Processing Systems*, 34:20172–20184, 2021.

- [16] Puning Zhao, Lifeng Lai, Li Shen, Qingming Li, Jiafei Wu, and Zhe Liu. A huber loss minimization approach to mean estimation under user-level differential privacy. *arXiv preprint arXiv:2405.13453*, 2024.
- [17] Hyun-Young Park, Shahab Asodeh, and Si-Hyeon Lee. Exactly minimax-optimal locally differentially private sampling. *Advances in Neural Information Processing Systems*, 37:10274–10319, 2024.
- [18] Kaan Ozkara, Antonious M. Girgis, Deepesh Data, and Suhas Diggavi. A statistical framework for personalized federated learning and estimation: Theory, algorithms, and privacy. In *The Eleventh International Conference on Learning Representations*, 2023. URL https://openreview.net/forum?id=FUidMCR_W4o.
- [19] Tian Li, Shengyuan Hu, Ahmad Beirami, and Virginia Smith. Ditto: Fair and robust federated learning through personalization. In Marina Meila and Tong Zhang, editors, *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pages 6357–6368. PMLR, 18–24 Jul 2021. URL <https://proceedings.mlr.press/v139/li21h.html>.
- [20] Michael Zhang, Karan Sapra, Sanja Fidler, Serena Yeung, and José M. Álvarez. Personalized federated learning with first order model optimization. In *ICLR*. OpenReview.net, 2021. URL <http://dblp.uni-trier.de/db/conf/iclr/iclr2021.html#ZhangSFYA21>.
- [21] Jinshuo Dong, Aaron Roth, and Weijie J Su. Gaussian differential privacy. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 84(1):3–37, 2022.
- [22] Behnoosh Zamanlooy, Mario Diaz, and Shahab Asodeh. Locally private sampling with public data. *arXiv preprint arXiv:2411.08791*, 2024.
- [23] Sofya Raskhodnikova, Satchit Sivakumar, Adam Smith, and Marika Swanberg. Differentially private sampling from distributions. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 28983–28994. Curran Associates, Inc., 2021. URL https://proceedings.neurips.cc/paper_files/paper/2021/file/f2b5e92f61b6de923b063588ee6e7c48-Paper.pdf.
- [24] Badih Ghazi, Xiao Hu, Ravi Kumar, and Pasin Manurangsi. On differentially private sampling from gaussian and product distributions. In A. Oh, T. Naumann, A. Globerson, K. Saenko, M. Hardt, and S. Levine, editors, *Advances in Neural Information Processing Systems*, volume 36, pages 77783–77809. Curran Associates, Inc., 2023. URL https://proceedings.neurips.cc/paper_files/paper/2023/file/f4eaa4b8f2d08edb3f0af990d56134ea-Paper-Conference.pdf.
- [25] Valentio Iverson, Gautam Kamath, and Argyris Mouzakis. Optimal differentially private sampling of unbounded gaussians. *arXiv preprint arXiv:2503.01766*, 2025.
- [26] Naima Tasnim, Atefeh Gilani, Lalitha Sankar, and Oliver Kosut. Reveal-or-obscure: A differentially private sampling algorithm for discrete distributions. *arXiv preprint arXiv:2504.14696*, 2025.
- [27] Albert Cheu and Debanuj Nayak. Differentially private multi-sampling from distributions, 2024. URL <https://arxiv.org/abs/2412.10512>.
- [28] James Flemings, Meisam Razaviyayn, and Murali Annavaram. Differentially private next-token prediction of large language models. *arXiv preprint arXiv:2403.15638*, 2024.
- [29] James Flemings, Meisam Razaviyayn, and Murali Annavaram. Adaptively private next-token prediction of large language models, 2024. URL <https://openreview.net/forum?id=fGSEWgRHNZ>.
- [30] Sofya Raskhodnikova, Satchit Sivakumar, Adam Smith, and Marika Swanberg. Differentially private sampling from distributions. *Advances in Neural Information Processing Systems*, 34: 28983–28994, 2021.

- [31] Badih Ghazi, Xiao Hu, Ravi Kumar, and Pasin Manurangsi. On differentially private sampling from gaussian and product distributions. *Advances in Neural Information Processing Systems*, 36:77783–77809, 2023.
- [32] Hamid Ebadi, Thibaud Antignac, and David Sands. Sampling and partitioning for differential privacy. In *2016 14th Annual Conference on Privacy, Security and Trust (PST)*, pages 664–673. IEEE, 2016.
- [33] Nazmiye Ceren Abay, Yan Zhou, Murat Kantarcioglu, Bhavani Thuraisingham, and Latanya Sweeney. Privacy preserving synthetic data release using deep learning. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pages 510–526. Springer, 2018.
- [34] Liyang Xie, Kaixiang Lin, Shu Wang, Fei Wang, and Jiayu Zhou. Differentially private generative adversarial network. *arXiv preprint arXiv:1802.06739*, 2018.
- [35] Bangzhou Xin, Wei Yang, Yangyang Geng, Sheng Chen, Shaowei Wang, and Liusheng Huang. Private fl-gan: Differential privacy synthetic data generation based on federated learning. In *Icassp 2020-2020 IEEE international conference on acoustics, speech and signal processing (ICASSP)*, pages 2927–2931. IEEE, 2020.
- [36] Reihaneh Torkzadehmahani, Peter Kairouz, and Benedict Paten. Dp-cgan: Differentially private synthetic data and label generation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops*, pages 0–0, 2019.
- [37] Yi Liu, Jialiang Peng, JQ James, and Yi Wu. Ppgan: Privacy-preserving generative adversarial network. In *2019 IEEE 25th international conference on parallel and distributed systems (ICPADS)*, pages 985–989. IEEE, 2019.
- [38] Teddy Cunningham, Konstantin Klemmer, Hongkai Wen, and Hakan Ferhatosmanoglu. Geopointgan: Synthetic spatial data with local label differential privacy. *arXiv preprint arXiv:2205.08886*, 2022.
- [39] Hisaichi Shibata, Shouhei Hanaoka, Yang Cao, Masatoshi Yoshikawa, Tomomi Takenaga, Yukihiko Nomura, Naoto Hayashi, and Osamu Abe. Local differential privacy image generation using flow-based deep generative models. *Applied Sciences*, 13(18):10132, 2023.
- [40] Hua Zhang, Kaixuan Li, Teng Huang, Xin Zhang, Wenmin Li, Zhengping Jin, Fei Gao, and Minghui Gao. Publishing locally private high-dimensional synthetic data efficiently. *Information Sciences*, 633:343–356, 2023.
- [41] Hansle Gwon, Imjin Ahn, Yunha Kim, Hee Jun Kang, Hyeram Seo, Heejung Choi, Ha Na Cho, Minkyung Kim, JiYe Han, Gaeun Kee, et al. Ldp-gan: generative adversarial networks with local differential privacy for patient medical records synthesis. *Computers in Biology and Medicine*, 168:107738, 2024.
- [42] Jacob Imola, Amrita Roy Chowdhury, and Kamalika Chaudhuri. Metric differential privacy at the user-level via the earth-mover’s distance. In *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security*, pages 348–362, 2024.
- [43] Wei-Ning Chen and Ayfer Özgür. Lq lower bounds on distributed estimation via fisher information. In *2024 IEEE International Symposium on Information Theory (ISIT)*, pages 91–96. IEEE, 2024.
- [44] Audra McMillan, Adam Smith, and Jon Ullman. Instance-optimal differentially private estimation. *arXiv preprint arXiv:2210.15819*, 2022.
- [45] John C. Duchi and Feng Ruan. The right complexity measure in locally private estimation: It is not the Fisher information. *The Annals of Statistics*, 52(1):1 – 51, 2024. doi: 10.1214/22-AOS2227. URL <https://doi.org/10.1214/22-AOS2227>.
- [46] Hilal Asi and John C Duchi. Instance-optimality in differential privacy via approximate inverse sensitivity mechanisms. *Advances in neural information processing systems*, 33:14106–14117, 2020.

- [47] Vitaly Feldman, Audra McMillan, Satchit Sivakumar, and Kunal Talwar. Instance-optimal private density estimation in the wasserstein distance. *Advances in Neural Information Processing Systems*, 37:90061–90131, 2024.
- [48] Angelika Rohde and Lukas Steinberger. Geometrizing rates of convergence under local differential privacy constraints. *The Annals of Statistics*, 48(5):2646–2670, 2020.
- [49] Larry Wasserman and Shuheng Zhou. A statistical framework for differential privacy. *Journal of the American Statistical Association*, 105(489):375–389, 2010.
- [50] Peter Kairouz, Sewoong Oh, and Pramod Viswanath. The composition theorem for differential privacy. In *International conference on machine learning*, pages 1376–1385. PMLR, 2015.
- [51] Bonwoo Lee, Jeongyoun Ahn, and Cheolwoo Park. Minimax risks and optimal procedures for estimation under functional local differential privacy. *Advances in Neural Information Processing Systems*, 36:57964–57975, 2023.
- [52] Syed Mumtaz Ali and Samuel D Silvey. A general class of coefficients of divergence of one distribution from another. *Journal of the Royal Statistical Society: Series B (Methodological)*, 28(1):131–142, 1966.
- [53] Imre Csiszár. On information-type measure of difference of probability distributions and indirect observations. *Studia Sci. Math. Hungar.*, 2:299–318, 1967.
- [54] Naresh Sharma and Naqeeb Ahmad Warsi. Fundamental bound on the reliability of quantum information transmission. *Physical review letters*, 110(8):080501, 2013.
- [55] Jingbo Liu, Paul Cuff, and Sergio Verdú. E_γ -resolvability. *IEEE Transactions on Information Theory*, 63(5):2629–2658, 2016.
- [56] Yury Polyanskiy and Yihong Wu. Dissipation of information in channels with input constraints. *IEEE Transactions on Information Theory*, 62(1):35–55, 2015.
- [57] Maxim Raginsky. Strong data processing inequalities and ϕ -sobolev inequalities for discrete channels. *IEEE Transactions on Information Theory*, 62(6):3355–3389, 2016.
- [58] Behnoosh Zamanlooy and Shahab Asoodeh. Strong data processing inequalities for locally differentially private mechanisms. In *2023 IEEE International Symposium on Information Theory (ISIT)*, pages 1794–1799. IEEE, 2023.
- [59] Shahab Asoodeh and Huanyu Zhang. Contraction of locally differentially private mechanisms. *IEEE Journal on Selected Areas in Information Theory*, 2024.
- [60] Igal Sason and Sergio Verdú. f -divergence inequalities. *IEEE Transactions on Information Theory*, 62(11):5973–6006, 2016.
- [61] Andrew Rukhin. Information-type divergence when the likelihood ratios are bounded. *Applicaciones Mathematicae*, 24(4):415–423, 1997.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The main contributions and scope of our paper are clearly outlined in the abstract and in introduction (Section 1).

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The limitations of our work is stated in Section 7.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: All theoretical results in the paper, including theorems, propositions, and lemmas explicitly state the necessary assumptions and are supported by either detailed proofs in the appendices or appropriate references.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: The detailed explanation of all of the experiments is provide in corresponding section or appendix. Moreover, the instructions for reproducing all figures and experimental results are provided in Appendix F.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: The complete code for generating all experimental results is made publicly available at https://github.com/hradghoukasian/private_sampling, and for completeness, the code is also included in the supplementary material. In addition, detailed instructions for reproducing all figures and experimental results are provided in Appendix F.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We provide sufficient details to implement our proposed sampler and include the code for both the proposed method and the baseline in the supplementary material.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: Since the main contribution lies in theoretically characterizing the *worst-case* utility, our experiments are designed to extract and evaluate this worst-case behavior, rather than to assess average-case performance or variability.

Guidelines:

- The answer NA means that the paper does not include experiments.

- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: Details on the computational resources and running times for the experiments are provided in Appendix F.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines?>

Answer: [Yes]

Justification: Our research conforms with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We discussed both the positive and negative societal impacts of our work in Appendix G.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: We do not identify any foreseeable risk of misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We cite the original paper associated with the source code and ensure that its license and terms of use are properly respected.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.

- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.

- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

Roadmap of Appendix: The appendix is organized as follows. A more complete notation table—complementing the abbreviated summary in the main body—is provided in Section A. Supplementary definitions and theorems used in the proofs of our results appear in Section B. Section C presents preliminary results that support the proofs of our main theorems, which are then detailed in Section D. Section E describes our experimental setup and provides additional numerical results. Finally, instructions for reproducing our experiments and a discussion of the broader impacts of our work are given in Sections F and G, respectively.

A Notation table

Table 1 provides a more comprehensive summary of the notation used throughout the paper, complementing the abbreviated overview in the main body for the reader’s convenience.

Table 1: Important notations used in the paper

Notation	Description
$\mathcal{P}(\mathcal{X})$	Set of all probability distributions over sample space $\mathcal{X}$
$\mathcal{C}(\mathbb{R}^n)$	Set of all continuous probability distributions on $\mathbb{R}^n$
p	Probability density function (PDF) of $P \in \mathcal{C}(\mathbb{R}^n)$
$[k]$	Shorthand for $\{1, 2, \dots, k\}$ for $k \in \mathbb{N}$
μ_k	uniform distribution over finite sample space $[k]$
$\text{clip}(x; s_1, s_2)$	Clipping function: $\max\{s_1, \min\{s_2, x\}\}$
g^*	Convex conjugate of a function g
$\mathcal{L}(m, b)$	n -dimensional Laplace distribution with mean $m \in \mathbb{R}^n$ and scale $b > 0$
$\mathcal{L}(x \mid m, b)$	Density of the Laplace distribution at $x \in \mathbb{R}^n$
$\mathcal{N}(m, \Sigma)$	n -dimensional Gaussian with mean $m \in \mathbb{R}^n$ and covariance matrix Σ
$\mathcal{N}_{m, \Sigma}[x]$	Density of the Gaussian distribution at $x \in \mathbb{R}^n$
$\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}$	Set of ε -LDP samplers: $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$
$\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}$	Set of (ε, δ) -LDP samplers: $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$
$\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}$	Set of g -FLDP samplers: $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$
$D_f^B(\lambda_1 \parallel \lambda_2)$	f -divergence between Bernoulli distributions with $\Pr(1) = \lambda_1$ and λ_2

B Supplementary definitions and theorems

B.1 General definition and special cases of f -divergence

In this appendix, we review the general definition and special cases of f -divergences which are important in this work.

Definition B.1 (General case of f -divergence [60]). Let P and Q be probability measures, and let μ be a dominating measure of P and Q (i.e., $P, Q \ll \mu$; e.g., $\mu = P + Q$), and let $p := \frac{dP}{d\mu}$ and $q := \frac{dQ}{d\mu}$. The f -divergence from P to Q is given, independently of μ , by

$$D_f(P \parallel Q) := \int q f\left(\frac{p}{q}\right) d\mu, \quad (10)$$

where

$$f(0) := \lim_{t \rightarrow 0^+} f(t), \quad (11)$$

$$0f\left(\frac{0}{0}\right) := 0, \quad (12)$$

$$0f\left(\frac{a}{0}\right) := \lim_{t \rightarrow 0^+} t f\left(\frac{a}{t}\right) = a \lim_{u \rightarrow \infty} \frac{f(u)}{u}, \quad a > 0. \quad (13)$$

Popular examples of f -divergences include KL divergence ($f(t) = t \log t$), total variation distance ($f(t) = \frac{1}{2}|t-1|$), squared Hellinger distance ($f(t) = (\sqrt{t}-1)^2$), and χ^2 -divergence ($f(t) = (t-1)^2$). We also formally define E_γ -divergence for $f(t) = \max\{t - \gamma, 0\}$ for $\gamma \geq 1$.

Definition B.2 (E_γ -divergence [54]). Given $\gamma \geq 1$, for two probability measures $P, Q \in \mathcal{P}(\mathcal{X})$ with $P \ll Q$, the E_γ -divergence is

$$E_\gamma(P \parallel Q) := \mathbb{E}_Q \left[\max \left\{ \frac{dP}{dQ} - \gamma, 0 \right\} \right] = \frac{1}{2} \int |dP - \gamma dQ| - \frac{1}{2}(\gamma - 1).$$

For a more extensive list of f -divergences and their properties, we refer readers to Sason and Verdú [60].

B.2 Measure-theoretic assumptions

The appendices assume familiarity with basic measure theory and real analysis. Throughout the main paper and the appendices, we adopt the following conventions. For every sample space $\mathcal{X}$, a σ -algebra on $\mathcal{X}$ is assumed to be given implicitly. Unless stated otherwise, we use the discrete σ -algebra when $\mathcal{X}$ is finite, and the Borel σ -algebra when $\mathcal{X} = \mathbb{R}^n$. Whenever we refer to a “subset” of $\mathcal{X}$, we mean a measurable subset. Similarly, the notation $A \subseteq \mathcal{X}$ always indicates that A is measurable. Finally, the term “continuous distribution” refers specifically to a distribution that is absolutely continuous with respect to the Lebesgue measure.

B.3 Supplementary theorems

This subsection presents two well-known results that we will use in proving our main results.

Theorem B.1 (Data processing inequality). *Let M be a conditional distribution (Markov kernel) mapping from $\mathcal{X}$ to $\mathcal{Y}$. Given distributions $P_1, P_2 \in \mathcal{P}(\mathcal{X})$, and their push-forward measures $Q_1, Q_2 \in \mathcal{P}(\mathcal{Y})$ through M , the following inequality holds for any f -divergence:*

$$D_f(P_1 \parallel P_2) \geq D_f(Q_1 \parallel Q_2).$$

Proposition B.2 (Proposition 1 of Dong et al. [21]). *A function $g : [0, 1] \rightarrow [0, 1]$ is a trade-off function if and only if g is convex, continuous, non-increasing, and $g(x) \leq 1 - x$ for $x \in [0, 1]$.*

C Preliminary results

Lemma C.1. *Let μ be a (fixed) distribution over some output space $\mathcal{X}$. For each $0 \leq \lambda \leq 1$, define a sampler $\mathbf{Q}_\lambda$ by*

$$\mathbf{Q}_\lambda(A \mid P) = \lambda P(A) + (1 - \lambda) \mu(A), \quad \text{for any distribution } P \text{ on } \mathcal{X} \text{ and event } A \subseteq \mathcal{X}.$$

Suppose there exists $0 \leq \lambda_1 \leq 1$ such that $\mathbf{Q}_{\lambda_1}$ is (ε, δ) -LDP. Then for any $\lambda_2 \in [0, \lambda_1]$, the sampler $\mathbf{Q}_{\lambda_2}$ is also (ε, δ) -LDP.

Proof. Fix $\lambda_2 \in [0, \lambda_1]$, and let P, P' be any two distributions on $\mathcal{X}$. We need to show that for every event A ,

$$\mathbf{Q}_{\lambda_2}(A \mid P) \leq e^\varepsilon \mathbf{Q}_{\lambda_2}(A \mid P') + \delta.$$

Because $\lambda_2 \leq \lambda_1$, there is a $\beta \in [0, 1]$ such that

$$\lambda_2 = \beta \lambda_1.$$

Then, for all events A ,

$$\mathbf{Q}_{\lambda_2}(A \mid P) = \lambda_2 P(A) + (1 - \lambda_2) \mu(A) = \beta \lambda_1 P(A) + [1 - \beta \lambda_1] \mu(A).$$

Recognizing that $\lambda_1 P(A) + (1 - \lambda_1) \mu(A)$ is $\mathbf{Q}_{\lambda_1}(A \mid P)$, we rewrite:

$$\mathbf{Q}_{\lambda_2}(A \mid P) = \beta \mathbf{Q}_{\lambda_1}(A \mid P) + [1 - \beta] \mu(A).$$

Similarly,

$$\mathbf{Q}_{\lambda_2}(A \mid P') = \beta \mathbf{Q}_{\lambda_1}(A \mid P') + [1 - \beta] \mu(A).$$

Since $\mathbf{Q}_{\lambda_1}$ is (ε, δ) -LDP, we know

$$\mathbf{Q}_{\lambda_1}(A \mid P) \leq e^\varepsilon \mathbf{Q}_{\lambda_1}(A \mid P') + \delta.$$

Multiplying both sides by $\beta \geq 0$ and then adding $(1 - \beta) \mu(A)$, we obtain

$$\beta \mathbf{Q}_{\lambda_1}(A | P) + (1 - \beta) \mu(A) \leq \beta e^\varepsilon \mathbf{Q}_{\lambda_1}(A | P') + \beta \delta + (1 - \beta) \mu(A).$$

But the left-hand side above is exactly $\mathbf{Q}_{\lambda_2}(A | P)$, so

$$\mathbf{Q}_{\lambda_2}(A | P) \leq \beta e^\varepsilon \mathbf{Q}_{\lambda_1}(A | P') + \beta \delta + (1 - \beta) \mu(A).$$

On the other hand,

$$e^\varepsilon \mathbf{Q}_{\lambda_2}(A | P') = e^\varepsilon [\beta \mathbf{Q}_{\lambda_1}(A | P') + (1 - \beta) \mu(A)] = \beta e^\varepsilon \mathbf{Q}_{\lambda_1}(A | P') + e^\varepsilon (1 - \beta) \mu(A).$$

Hence,

$$e^\varepsilon \mathbf{Q}_{\lambda_2}(A | P') + \delta = \beta e^\varepsilon \mathbf{Q}_{\lambda_1}(A | P') + (1 - \beta) e^\varepsilon \mu(A) + \delta.$$

It is enough to check

$$\beta e^\varepsilon \mathbf{Q}_{\lambda_1}(A | P') + \beta \delta + (1 - \beta) \mu(A) \leq \beta e^\varepsilon \mathbf{Q}_{\lambda_1}(A | P') + (1 - \beta) e^\varepsilon \mu(A) + \delta,$$

which simplifies to

$$\beta \delta + (1 - \beta) \mu(A) \leq (1 - \beta) e^\varepsilon \mu(A) + \delta.$$

Rearranging,

$$(1 - \beta) \mu(A) [1 - e^\varepsilon] \leq \delta [1 - \beta].$$

Since $1 - e^\varepsilon \leq 0$ for $\varepsilon \geq 0$, and $\delta \geq 0$, this inequality holds trivially (the left-hand side is at most 0, while the right-hand side is nonnegative). Thus the overall chain of inequalities is valid, and we conclude

$$\mathbf{Q}_{\lambda_2}(A | P) \leq e^\varepsilon \mathbf{Q}_{\lambda_2}(A | P') + \delta, \quad \forall P, P', \text{ event } A.$$

Hence $\mathbf{Q}_{\lambda_2}$ is also (ε, δ) -LDP. $\square$

Now, we first restate the normalization condition and introduce the non-triviality assumption for (ε, δ) -LDP before stating the next proposition.

Normalization condition:

$$\mu(\mathcal{X}) = 1, \quad 0 \leq c_1 < 1 < c_2. \quad (14)$$

Non-triviality condition:

$$\delta \leq \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1}. \quad (15)$$

Proposition C.2. Consider the following optimization problem:

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}, \tilde{\mathcal{P}}, f) := \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)). \quad (16)$$

Let $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$ such that the normalization condition (14) and the non-triviality condition (15) hold and $\frac{c_2 - c_1}{1 - c_1} \in \mathbb{N}$. Suppose that μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1 - c_1}{c_2 - c_1}$. Define

$$r_1 := \frac{c_1}{c_2 - c_1} \cdot \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{1 - \delta}, \quad r_2 := \frac{c_2}{c_2 - c_1} \cdot \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{e^\varepsilon + \frac{c_2 - 1}{1 - c_1} \delta}.$$

Then, the optimal value of the problem (16) is given by

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}, \tilde{\mathcal{P}}, f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1). \quad (17)$$

Moreover, the sampler $\mathbf{Q}_{\varepsilon, \delta}^*(P) := \lambda_{\varepsilon, \delta}^* P + (1 - \lambda_{\varepsilon, \delta}^*) \mu$, with $\lambda_{\varepsilon, \delta}^* = \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} \delta - 1}{(1 - c_1)e^\varepsilon + c_2 - 1}$, belongs to the class $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}$ and is minimax-optimal under any f -divergence D_f . That is,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{\varepsilon, \delta}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}, \tilde{\mathcal{P}}, f). \quad (18)$$

Proof. We first show that if the non-triviality constraint (15) does not hold, then the identity sampler $\mathbf{Q}^I(P) = P$ satisfies (ε, δ) -LDP and has the trivial minimax risk zero.

Suppose we have $\delta > \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1}$. Let $A \subseteq \mathcal{X}$ be a measurable subset and $P, P' \in \mathcal{P}_{c_1, c_2, \mu}$. If $0 \leq \mu(A) \leq \frac{1 - c_1}{c_2 - c_1}$,

$$\begin{aligned} e^\varepsilon P(A) + \delta - P'(A) &\geq c_1 e^\varepsilon \mu(A) + \delta - c_2 \mu(A) \\ &\geq \min \left\{ \frac{(c_1 e^\varepsilon - c_2)(1 - c_1)}{c_2 - c_1} + \delta, \delta \right\} \\ &\geq 0. \end{aligned}$$

If $\mu(A) \geq \frac{1 - c_1}{c_2 - c_1}$,

$$\begin{aligned} e^\varepsilon P(A) + \delta - P'(A) &\geq c_1 e^\varepsilon \mu(A) + \delta - (1 - P'(A^c)) \\ &\geq c_1 e^\varepsilon \mu(A) + \delta - (1 - c_1(1 - \mu(A))) \\ &= (c_1 e^\varepsilon - c_1) \mu(A) + \delta - 1 + c_1 \\ &\geq \frac{(c_1 e^\varepsilon - c_1)(1 - c_1)}{c_2 - c_1} - 1 + c_1 + \delta \\ &= \frac{(c_1 e^\varepsilon - c_2)(1 - c_1)}{c_2 - c_1} + \delta \\ &> 0. \end{aligned}$$

Therefore, for $\delta > \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1}$, $\mathbf{Q}^I(P) \in \mathcal{Q}_{\mathcal{X}, \tilde{P}, \varepsilon, \delta}$ and the minimax risk is trivially zero. In the remainder of the proof, we consider the non-trivial case of $\delta \leq \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1}$.

Step 1: Proof of $\mathbf{Q}_{\varepsilon, \delta}^*$ is a probability measure

In this proof, we show that $\mathbf{Q}_{\varepsilon, \delta}^*$ defined by $\mathbf{Q}_{\varepsilon, \delta}^*(P) := \lambda_{\varepsilon, \delta}^* P + (1 - \lambda_{\varepsilon, \delta}^*) \mu$ is a probability measure. Since P and μ are probability measures, we have

$$\mathbf{Q}_{\varepsilon, \delta}^*(\mathcal{X}) = \lambda_{\varepsilon, \delta}^* P(\mathcal{X}) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(\mathcal{X}) = \lambda_{\varepsilon, \delta}^* + (1 - \lambda_{\varepsilon, \delta}^*) = 1.$$

Assumptions $0 \leq c_1 < 1 < c_2$ and $\delta \leq \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1}$ imply

$$\lambda_{\varepsilon, \delta}^* = \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} \delta - 1}{(1 - c_1) e^\varepsilon + c_2 - 1} \leq \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1} - 1}{(1 - c_1) e^\varepsilon + c_2 - 1} = 1.$$

Therefore, for any $A \subseteq \mathcal{X}$, we have $\mathbf{Q}_{\varepsilon, \delta}^*(A) \geq 0$. It suffices to show that for every measurable subset $A \subseteq \mathcal{X}$, we also have

$$\mathbf{Q}_{\varepsilon, \delta}^*(A) \leq 1,$$

which implies $\mathbf{Q}_{\varepsilon, \delta}^*$ is a probability measure. We proceed by bounding:

$$\begin{aligned} \mathbf{Q}_{\varepsilon, \delta}^*(A) &= \lambda_{\varepsilon, \delta}^* P(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A) \\ &\leq \lambda_{\varepsilon, \delta}^* (1 - c_1(1 - \mu(A))) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A) \\ &= (1 - (1 - c_1) \lambda_{\varepsilon, \delta}^*) \mu(A) + \lambda_{\varepsilon, \delta}^* (1 - c_1) \\ &\leq \frac{(c_2 - c_1)(1 - \delta)}{(1 - c_1) e^\varepsilon + c_2 - 1} + \frac{(1 - c_1) e^\varepsilon + (c_2 - c_1) \delta - (1 - c_1)}{(1 - c_1) e^\varepsilon + c_2 - 1} \\ &= \frac{(1 - c_1) e^\varepsilon + c_2 - 1}{(1 - c_1) e^\varepsilon + c_2 - 1} \\ &= 1. \end{aligned}$$

The inequality in the second line uses the fact that

$$P(A) \leq 1 - c_1(1 - \mu(A)).$$

This follows from the definition

$$\tilde{\mathcal{P}}_{c_1, c_2, \mu} := \left\{ P \in \mathcal{P}(\mathcal{X}) : P \ll \mu, \quad c_1 \leq \frac{dP}{d\mu} \leq c_2 \quad \mu\text{-a.e.} \right\},$$

and the fact that for $P \in \tilde{\mathcal{P}}_{c_1, c_2, \mu}$:

$$c_1 \mu(A^c) \leq P(A^c) \implies 1 - c_1 \mu(A^c) \geq 1 - P(A^c) \implies 1 - c_1(1 - \mu(A)) \geq P(A).$$

And the inequality in the fourth line comes by plugging $\lambda_{\varepsilon, \delta}^* = \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} \delta - 1}{(1 - c_1)e^\varepsilon + c_2 - 1}$ into the equation and the fact that $\mu(A) \leq 1$. Hence, $\mathbf{Q}_{\varepsilon, \delta}^*$ is indeed a probability measure.

Step 2: Proof of $\mathbf{Q}_{\varepsilon, \delta}^* \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}$

Now, we want to show that $\mathbf{Q}_{\varepsilon, \delta}^*$ is an (ε, δ) -LDP sampler. For this purpose, we have to show that for each measurable set $A \subseteq \mathcal{X}$ and any pair of probability measures P and P' , we have:

$$e^\varepsilon (\mathbf{Q}_{\varepsilon, \delta}^*(P)(A)) + \delta - \mathbf{Q}_{\varepsilon, \delta}^*(P')(A) \geq 0.$$

First, suppose $\mu(A) \leq \frac{1 - c_1}{c_2 - c_1}$. We then have:

$$\begin{aligned} & e^\varepsilon (\lambda_{\varepsilon, \delta}^* P(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) + \delta - (\lambda_{\varepsilon, \delta}^* P'(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) \\ & \geq e^\varepsilon (\lambda_{\varepsilon, \delta}^* c_1 \mu(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) + \delta - (\lambda_{\varepsilon, \delta}^* c_2 \mu(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) \\ & = (\lambda_{\varepsilon, \delta}^* (e^\varepsilon c_1 - e^\varepsilon - c_2 + 1) + (e^\varepsilon - 1)) \mu(A) + \delta \\ & = -\frac{c_2 - c_1}{1 - c_1} \delta \mu(A) + \delta \\ & \geq 0. \end{aligned}$$

Next, suppose $\mu(A) \geq \frac{1 - c_1}{c_2 - c_1}$. Then:

$$\begin{aligned} & e^\varepsilon (\lambda_{\varepsilon, \delta}^* P(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) + \delta - (\lambda_{\varepsilon, \delta}^* P'(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) \\ & \geq e^\varepsilon (\lambda_{\varepsilon, \delta}^* c_1 \mu(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) + \delta - (\lambda_{\varepsilon, \delta}^* (1 - c_1 + c_1 \mu(A)) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) \\ & = (e^\varepsilon (1 - (1 - c_1) \lambda_{\varepsilon, \delta}^*) - (1 - (1 - c_1) \lambda_{\varepsilon, \delta}^*)) \mu(A) - \lambda_{\varepsilon, \delta}^* (1 - c_1) + \delta \\ & = (e^\varepsilon - 1) (1 - (1 - c_1) \lambda_{\varepsilon, \delta}^*) \mu(A) - \lambda_{\varepsilon, \delta}^* (1 - c_1) + \delta \\ & \geq (e^\varepsilon - 1) (1 - (1 - c_1) \lambda_{\varepsilon, \delta}^*) \left(\frac{1 - c_1}{c_2 - c_1} \right) - \lambda_{\varepsilon, \delta}^* (1 - c_1) + \delta \\ & = \frac{(e^\varepsilon - 1)(1 - \delta)(1 - c_1)}{(1 - c_1)e^\varepsilon + c_2 - 1} - \frac{(1 - c_1)e^\varepsilon + (c_2 - c_1)\delta - (1 - c_1)}{(1 - c_1)e^\varepsilon + c_2 - 1} + \delta \\ & = \frac{-\delta(1 - c_1)(e^\varepsilon - 1) - (c_2 - c_1)\delta}{(1 - c_1)e^\varepsilon + c_2 - 1} + \delta \\ & = 0. \end{aligned}$$

Therefore, in both cases:

$$e^\varepsilon (\lambda_{\varepsilon, \delta}^* P(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) + \delta - (\lambda_{\varepsilon, \delta}^* P'(A) + (1 - \lambda_{\varepsilon, \delta}^*) \mu(A)) \geq 0,$$

which establishes the (ε, δ) -LDP criterion for every measurable set A . Hence the sampler is (ε, δ) -LDP.

Step 3: Proof of optimality, lower bound

In this part, we show that:

$$\frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1) \leq \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}, \tilde{\mathcal{P}}, f).$$

To this end, we need to show that for each $\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}$, we have

$$\frac{1-r_1}{r_2-r_1} f(r_2) + \frac{r_2-1}{r_2-r_1} f(r_1) \leq \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P))$$

When $\frac{c_2-c_1}{1-c_1} \in \mathbb{N}$, let $\mathbf{Q}$ be an (ε, δ) -LDP sampler, i.e. $\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}$. Define

$$\alpha := \frac{1-c_1}{c_2-c_1}, \quad t := \alpha^{-1} \in \mathbb{N}.$$

Since μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\frac{1}{\alpha} \in \mathbb{N}$, we can find disjoint sets $A_1, \dots, A_t \subseteq \mathcal{X}$ with $\mu(A_i) = \alpha$. Define P_i as a probability measure by $\frac{dP_i}{d\mu} = p_i$ defined as:

$$p_i(x) = c_2 \mathbf{1}_{A_i}(x) + c_1 \mathbf{1}_{A_i^c}(x),$$

where

$$\mathbf{1}_{A_i}(x) = \begin{cases} 1 & \text{if } x \in A_i, \\ 0 & \text{otherwise.} \end{cases}$$

Let $\mathbf{Q}_i := \mathbf{Q}(P_i)$. We aim to show:

$$\min_i \mathbf{Q}_i(A_i) \leq \frac{e^\varepsilon + (t-1)\delta}{e^\varepsilon + t - 1}.$$

By the (ε, δ) -LDP property, for every $i > 1$, we have

$$e^\varepsilon \mathbf{Q}_1(A_i) + \delta \geq \mathbf{Q}_i(A_i).$$

Thus,

$$\sum_{i=2}^t (e^\varepsilon \mathbf{Q}_1(A_i) + \delta) \geq \sum_{i=2}^t \mathbf{Q}_i(A_i).$$

Note that

$$\begin{aligned} \sum_{i=2}^t (e^\varepsilon \mathbf{Q}_1(A_i) + \delta) &= e^\varepsilon \mathbf{Q}_1\left(\bigcup_{i=2}^t A_i\right) + (t-1)\delta \\ &\leq e^\varepsilon (1 - \mathbf{Q}_1(A_1)) + (t-1)\delta \\ &\leq e^\varepsilon \left(1 - \min_i \mathbf{Q}_i(A_i)\right) + (t-1)\delta, \end{aligned}$$

and

$$\sum_{i=2}^t \mathbf{Q}_i(A_i) \geq (t-1) \min_i \mathbf{Q}_i(A_i).$$

Therefore,

$$\min_i \mathbf{Q}_i(A_i) \leq \frac{e^\varepsilon + (t-1)\delta}{e^\varepsilon + t - 1}.$$

Rewriting in terms of α , we have

$$\min_i \mathbf{Q}_i(A_i) \leq \frac{\alpha e^\varepsilon + (1-\alpha)\delta}{\alpha e^\varepsilon + 1 - \alpha}. \quad (19)$$

Next, we lower bound the worst-case f -divergence:

$$\sup_i D_f(P_i \| \mathbf{Q}(P_i)) \geq \sup_i D_f^B(c_2\alpha \| \mathbf{Q}_i(A_i)),$$

where $D_f^B(\lambda_1 \| \lambda_2)$ denotes the f -divergence between Bernoulli distributions with $\Pr(1) = \lambda_1$ and λ_2 :

$$D_f^B(\lambda_1 \| \lambda_2) = \lambda_2 f\left(\frac{\lambda_1}{\lambda_2}\right) + (1-\lambda_2) f\left(\frac{1-\lambda_1}{1-\lambda_2}\right).$$

For each A_i , the push-forward measures of P_i and $\mathbf{Q}_i$ by the indicator function $\mathbf{1}_{A_i}$ are Bernoulli distributions with $\Pr(1) = c_2\alpha$ and $\mathbf{Q}_i(A_i)$, respectively. By the data processing inequality (Theorem B.1), we have:

$$D_f(P_i \| \mathbf{Q}_i) \geq D_f^B(c_2\alpha \| \mathbf{Q}_i(A_i)).$$

Therefore, we have:

$$\begin{aligned} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)) &\geq \sup_{i \in [t]} D_f(P_i \| \mathbf{Q}_i) \\ &\geq \sup_{i \in [t]} D_f^B(c_2\alpha \| \mathbf{Q}_i(A_i)) \\ &\geq D_f^B\left(c_2\alpha \| \frac{\alpha e^\varepsilon + (1-\alpha)\delta}{\alpha e^\varepsilon + 1 - \alpha}\right) \end{aligned}$$

where the inequality in the last line follows from (19) and the fact that if $\delta \leq \frac{(c_2 - c_1)e^\varepsilon(1 - c_1)}{c_2 - c_1}$, then $\frac{\alpha e^\varepsilon + (1-\alpha)\delta}{\alpha e^\varepsilon + 1 - \alpha} \leq c_2\alpha$ and $D_f^B(\lambda_1 \| \lambda_2)$ is decreasing in $\lambda_2 \in [0, \lambda_1]$.

Here is the proof that if $\delta \leq \frac{(c_2 - c_1)e^\varepsilon(1 - c_1)}{c_2 - c_1}$, then $\frac{\alpha e^\varepsilon + (1-\alpha)\delta}{\alpha e^\varepsilon + 1 - \alpha} \leq c_2\alpha$.

$$\begin{aligned} \frac{\alpha e^\varepsilon + (1-\alpha)\delta}{\alpha e^\varepsilon + 1 - \alpha} &= \frac{(1 - c_1)e^\varepsilon + (c_2 - 1)\delta}{(1 - c_1)e^\varepsilon + c_2 - 1} \\ &\leq \frac{(1 - c_1)e^\varepsilon + (c_2 - 1)\left(\frac{(c_2 - c_1)e^\varepsilon(1 - c_1)}{c_2 - c_1}\right)}{(1 - c_1)e^\varepsilon + c_2 - 1} \\ &= \frac{1 - c_1}{c_2 - c_1} \left(\frac{(c_2 - c_1)e^\varepsilon + (c_2 - 1)(c_2 - c_1 e^\varepsilon)}{(1 - c_1)e^\varepsilon + c_2 - 1} \right) \\ &= \frac{c_2(1 - c_1)}{c_2 - c_1} \\ &= c_2\alpha. \end{aligned}$$

Therefore, it suffices to compute $D_f^B\left(c_2\alpha \| \frac{\alpha e^\varepsilon + (1-\alpha)\delta}{\alpha e^\varepsilon + 1 - \alpha}\right)$. It can be shown that

$$\begin{aligned} D_f^B\left(c_2\alpha \| \frac{\alpha e^\varepsilon + (1-\alpha)\delta}{\alpha e^\varepsilon + 1 - \alpha}\right) &= \frac{\alpha e^\varepsilon + (1-\alpha)\delta}{\alpha e^\varepsilon + 1 - \alpha} f\left(c_2\alpha \frac{\alpha e^\varepsilon + 1 - \alpha}{\alpha e^\varepsilon + (1-\alpha)\delta}\right) \\ &\quad + \frac{(1-\alpha)(1-\delta)}{\alpha e^\varepsilon + 1 - \alpha} f\left(c_1(1-\alpha) \frac{\alpha e^\varepsilon + 1 - \alpha}{(1-\alpha)(1-\delta)}\right) \\ &= \frac{(1 - c_1)e^\varepsilon + (c_2 - 1)\delta}{(1 - c_1)e^\varepsilon + c_2 - 1} f\left(\frac{c_2}{c_2 - c_1} \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{e^\varepsilon + \frac{c_2 - 1}{1 - c_1}\delta}\right) \\ &\quad + \frac{(c_2 - 1)(1 - \delta)}{(1 - c_1)e^\varepsilon + c_2 - 1} f\left(\frac{c_1}{c_2 - c_1} \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{1 - \delta}\right). \end{aligned}$$

Defining

$$r_1 := \frac{c_1}{c_2 - c_1} \cdot \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{1 - \delta}, \quad r_2 := \frac{c_2}{c_2 - c_1} \cdot \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{e^\varepsilon + \frac{c_2 - 1}{1 - c_1}\delta},$$

we have

$$\begin{aligned} 1 - r_1 &= \frac{(1 - c_1)\left(c_2 - c_1 e^\varepsilon - \frac{c_2 - c_1}{1 - c_1}\delta\right)}{(c_2 - c_1)(1 - \delta)}, \\ r_2 - 1 &= \frac{(c_2 - 1)\left(c_2 - c_1 e^\varepsilon - \frac{c_2 - c_1}{1 - c_1}\delta\right)}{(c_2 - c_1)\left(e^\varepsilon + \frac{c_2 - 1}{1 - c_1}\delta\right)}. \end{aligned}$$

Thus

$$\begin{aligned}\frac{1-r_1}{r_2-r_1} &= \frac{(1-c_1)e^\varepsilon + (c_2-1)\delta}{(1-c_1)e^\varepsilon + c_2-1}, \\ \frac{r_2-1}{r_2-r_1} &= \frac{(c_2-1)(1-\delta)}{(1-c_1)e^\varepsilon + c_2-1}.\end{aligned}$$

Hence, for each $\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}$, we have

$$\frac{1-r_1}{r_2-r_1} f(r_2) + \frac{r_2-1}{r_2-r_1} f(r_1) \leq \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)).$$

Step 4: Proof of optimality, upper bound (achievability part)

In this part, we show that:

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}, \tilde{\mathcal{P}}, f) \leq \frac{1-r_1}{r_2-r_1} f(r_2) + \frac{r_2-1}{r_2-r_1} f(r_1).$$

To this end, we need to show that for each $P \in \tilde{\mathcal{P}}$, we have:

$$D_f(P \| \mathbf{Q}_{\varepsilon, \delta}^*(P)) \leq \frac{1-r_1}{r_2-r_1} f(r_2) + \frac{r_2-1}{r_2-r_1} f(r_1).$$

Fix $P \in \tilde{\mathcal{P}}$. Let $p = \frac{dP}{d\mu}$ and $q = \frac{d\mathbf{Q}_{\varepsilon, \delta}^*(P)}{d\mu}$. We first claim that

$$r_1 \leq \frac{p(x)}{q(x)} \leq r_2$$

for μ -almost every $x \in \mathcal{X}$. Indeed, we have

$$\frac{p(x)}{q(x)} = \frac{p(x)}{\lambda_{\varepsilon, \delta}^* p(x) + (1 - \lambda_{\varepsilon, \delta}^*)},$$

which is an increasing function of $p(x) \geq 0$ as long as $\lambda_{\varepsilon, \delta}^* \leq 1$. Now we show that if $\delta \leq \frac{(c_2-c_1)e^\varepsilon(1-c_1)}{c_2-c_1}$, then $\lambda_{\varepsilon, \delta}^* = \frac{e^\varepsilon + \frac{c_2-c_1}{1-c_1}\delta - 1}{(1-c_1)e^\varepsilon + c_2 - 1} \leq 1$.

$$\begin{aligned}\lambda_{\varepsilon, \delta}^* &= \frac{e^\varepsilon + \frac{c_2-c_1}{1-c_1}\delta - 1}{(1-c_1)e^\varepsilon + c_2 - 1} \\ &\leq \frac{e^\varepsilon + \frac{c_2-c_1}{1-c_1} \left(\frac{(c_2-c_1)e^\varepsilon(1-c_1)}{c_2-c_1} \right) - 1}{(1-c_1)e^\varepsilon + c_2 - 1} \\ &= \frac{e^\varepsilon + c_2 - c_1e^\varepsilon - 1}{(1-c_1)e^\varepsilon + c_2 - 1} \\ &= 1.\end{aligned}$$

Since $p(x)$ is bounded between c_1 and c_2 , it follows that

$$\frac{c_1}{\lambda_{\varepsilon, \delta}^* c_1 + (1 - \lambda_{\varepsilon, \delta}^*)} \leq \frac{p(x)}{q(x)} \leq \frac{c_2}{\lambda_{\varepsilon, \delta}^* c_2 + (1 - \lambda_{\varepsilon, \delta}^*)} \quad (20)$$

for μ -almost every $x \in \mathcal{X}$. From the premise of the proposition, we have

$$\lambda_{\varepsilon, \delta}^* = \frac{e^\varepsilon + \frac{c_2-c_1}{1-c_1}\delta - 1}{(1-c_1)e^\varepsilon + c_2 - 1}$$

It can be shown that:

$$\frac{c_1}{c_2 - c_1} \frac{(1-c_1)e^\varepsilon + c_2 - 1}{(1-\delta)} \leq \frac{p(x)}{q(x)} \leq \frac{c_2}{c_2 - c_1} \cdot \frac{(1-c_1)e^\varepsilon + c_2 - 1}{e^\varepsilon + \frac{c_2-1}{1-c_1}\delta}$$

which concludes

$$r_1 \leq \frac{p(x)}{q(x)} \leq r_2$$

for μ -almost every $x \in \mathcal{X}$. To finalize the proof of the achievability part, we need the following lemma.

Lemma C.3 ([Theorem 2.1][61]). Let $P, Q \in \mathcal{P}(\mathcal{X})$. Suppose that P and Q are both absolutely continuous with respect to a reference measure μ on $\mathcal{X}$. Assume that there exist $r_1, r_2 \in \mathbb{R}$ with $0 \leq r_1 < 1 < r_2$ such that the corresponding densities $p = \frac{dP}{d\mu}$ and $q = \frac{dQ}{d\mu}$ satisfy $q(x) > 0$ and

$$r_1 \leq \frac{p(x)}{q(x)} \leq r_2 \quad \text{for } \mu\text{-almost every } x \in \mathcal{X}.$$

Then, for any f -divergence D_f , it holds that

$$D_f(P \| Q) \leq \frac{1-r_1}{r_2-r_1} f(r_2) + \frac{r_2-1}{r_2-r_1} f(r_1).$$

Lemma C.3, directly shows that for each $P \in \tilde{\mathcal{P}}$, we have:

$$D_f(P \| \mathbf{Q}_{\varepsilon, \delta}^*(P)) \leq \frac{1-r_1}{r_2-r_1} f(r_2) + \frac{r_2-1}{r_2-r_1} f(r_1)$$

for

$$r_1 = \frac{c_1}{c_2 - c_1} \cdot \frac{(1-c_1)e^\varepsilon + c_2 - 1}{1 - \delta}, \quad r_2 = \frac{c_2}{c_2 - c_1} \cdot \frac{(1-c_1)e^\varepsilon + c_2 - 1}{e^\varepsilon + \frac{c_2-1}{1-c_1}\delta}.$$

□

Corollary C.4 ($g_{\varepsilon, \delta}$ -FLDP as a special case). Let $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$. Under Assumption 3.2, suppose the reference measure μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1-c_1}{c_2-c_1}$. Then the sampler

$$\mathbf{Q}_{c_1, c_2, \mu, g_{\varepsilon, \delta}}^*(P) := \lambda_{c_1, c_2, g_{\varepsilon, \delta}}^* P + (1 - \lambda_{c_1, c_2, g_{\varepsilon, \delta}}^*) \mu \quad (21)$$

belongs to the class $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g_{\varepsilon, \delta}}$ and is minimax-optimal with respect to any f -divergence D_f , that is,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, \mu, g_{\varepsilon, \delta}}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g_{\varepsilon, \delta}}, \tilde{\mathcal{P}}, f), \quad (22)$$

where

$$\lambda_{c_1, c_2, g_{\varepsilon, \delta}}^* = \frac{e^\varepsilon + \frac{c_2-c_1}{1-c_1} \delta - 1}{(1-c_1)e^\varepsilon + c_2 - 1}.$$

Proof. From Theorem 3.6, we know that the sampler

$$\mathbf{Q}_{c_1, c_2, \mu, g_{\varepsilon, \delta}}^*(P) := \lambda_{c_1, c_2, g_{\varepsilon, \delta}}^* P + (1 - \lambda_{c_1, c_2, g_{\varepsilon, \delta}}^*) \mu$$

for $\lambda_{c_1, c_2, g_{\varepsilon, \delta}}^*$ defined as

$$\lambda_{c_1, c_2, g_{\varepsilon, \delta}}^* := \inf_{\beta \geq 0} \frac{e^\beta + \frac{c_2-c_1}{1-c_1} (1 + g_{\varepsilon, \delta}^*(-e^\beta)) - 1}{(1-c_1)e^\beta + c_2 - 1}.$$

belongs to the class $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g_{\varepsilon, \delta}}$ and is minimax-optimal with respect to any f -divergence D_f . Therefore, it suffices to compute $g_{\varepsilon, \delta}^*(-e^\varepsilon)$ and then solve the minimization problem.

A direct comparison of the three affine pieces shows

$$g_{\varepsilon, \delta}(\theta) = \begin{cases} 1 - \delta - e^\varepsilon \theta, & 0 \leq \theta \leq \frac{1-\delta}{e^\varepsilon+1}, \\ e^{-\varepsilon}(1 - \delta - \theta), & \frac{1-\delta}{e^\varepsilon+1} \leq \theta \leq 1 - \delta, \\ 0, & 1 - \delta \leq \theta \leq 1, \\ +\infty, & \text{otherwise.} \end{cases} \quad (23)$$

For any $y \in \mathbb{R}$, $g_{\varepsilon, \delta}^*(y) = \sup_{0 \leq \theta \leq 1} (\theta y - g_{\varepsilon, \delta}(\theta))$ splits naturally into the supremum over the three intervals of (23). Denote the corresponding optimized values by

$$S_1(y) := \sup_{0 \leq \theta \leq \frac{1-\delta}{e^\varepsilon+1}} \{\theta y - (1 - \delta - e^\varepsilon \theta)\},$$

$$S_2(y) := \sup_{\frac{1-\delta}{e^\varepsilon+1} \leq \theta \leq 1-\delta} \{\theta y - e^{-\varepsilon}(1 - \delta - \theta)\},$$

$$S_3(y) := \sup_{1-\delta \leq \theta \leq 1} \{\theta y\}.$$

We now optimize over each sub-interval.

(i) Interval $0 \leq \theta \leq \frac{1-\delta}{e^\varepsilon+1}$. Writing the objective as $(y + e^\varepsilon)\theta - (1 - \delta)$, it is linear in θ . Hence

$$S_1(y) = \begin{cases} -(1 - \delta), & y \leq -e^\varepsilon, \\ (y + e^\varepsilon) \frac{1-\delta}{e^\varepsilon+1} - (1 - \delta) = \frac{1-\delta}{e^\varepsilon+1}(y - 1), & y \geq -e^\varepsilon. \end{cases}$$

(ii) Interval $\frac{1-\delta}{e^\varepsilon+1} \leq \theta \leq 1 - \delta$. Here the objective equals $(y + e^{-\varepsilon})\theta - e^{-\varepsilon}(1 - \delta)$. Thus

$$S_2(y) = \begin{cases} \frac{1-\delta}{e^\varepsilon+1}(y - 1), & y \leq -e^{-\varepsilon}, \\ (1 - \delta)y, & y \geq -e^{-\varepsilon}. \end{cases}$$

(iii) Interval $1 - \delta \leq \theta \leq 1$. The objective is simply $y\theta$, so

$$S_3(y) = \begin{cases} (1 - \delta)y, & y < 0, \\ y, & y \geq 0. \end{cases}$$

We now take the overall supremum.

Comparing S_1, S_2, S_3 on the four regimes $(-\infty, -e^\varepsilon)$, $[-e^\varepsilon, -e^{-\varepsilon}]$, $[-e^{-\varepsilon}, 0]$, $[0, \infty)$ gives

$$g_{\varepsilon, \delta}^*(y) = \max\{S_1(y), S_2(y), S_3(y)\} = \begin{cases} -(1 - \delta), & y < -e^\varepsilon, \\ \frac{1-\delta}{e^\varepsilon+1}(y - 1), & -e^\varepsilon \leq y \leq -e^{-\varepsilon}, \\ (1 - \delta)y, & -e^{-\varepsilon} \leq y \leq 0, \\ y, & y \geq 0. \end{cases}$$

Step 2: Solve the infimum to obtain $\lambda_{c_1, c_2, g_{\varepsilon, \delta}}^*$

From

$$\lambda_{c_1, c_2, g_{\varepsilon, \delta}}^* = \inf_{\beta \geq 0} \frac{e^\beta + \frac{c_2 - c_1}{1 - c_1} (1 + g_{\varepsilon, \delta}^*(-e^\beta)) - 1}{(1 - c_1)e^\beta + c_2 - 1}$$

Hence, we only need to know the value of $g_{\varepsilon, \delta}^*(-e^\beta)$. We know $-e^\beta \leq -1$ for $\beta \geq 0$. Therefore, $-e^\beta \leq -e^{-\varepsilon}$ for given $\varepsilon \geq 0$. i.e.,

$$g_{\varepsilon, \delta}^*(-e^\beta) = \begin{cases} -(1 - \delta), & -e^\beta \leq -e^\varepsilon, \\ \frac{1-\delta}{e^\varepsilon+1}(-e^\beta - 1), & -e^\varepsilon \leq -e^\beta \leq -1. \end{cases}$$

It follows that:

$$\begin{aligned} \lambda_{c_1, c_2, g_{\varepsilon, \delta}}^* &= \min \left\{ \inf_{\beta \geq \varepsilon} \frac{e^\beta + \frac{c_2 - c_1}{1 - c_1} \delta - 1}{(1 - c_1)e^\beta + c_2 - 1}, \inf_{\beta \in [0, \varepsilon]} \frac{e^\beta + \frac{c_2 - c_1}{1 - c_1} \left(1 - \frac{e^\beta + 1}{e^\varepsilon + 1}(1 - \delta)\right) - 1}{(1 - c_1)e^\beta + c_2 - 1} \right\} \\ &= \min \left\{ \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} \delta - 1}{(1 - c_1)e^\varepsilon + c_2 - 1}, \frac{1 - (1 - \delta) \frac{2}{e^\varepsilon + 1}}{1 - c_1} \right\}. \end{aligned}$$

Since $\frac{c_2 - c_1}{1 - c_1} \in \mathbb{N}$ and $c_2 > 1$, it follows that $\frac{c_2 - c_1}{1 - c_1} \geq 2$, which in turn implies $c_1 + c_2 \geq 2$. Now we want to show that if $c_1 + c_2 \geq 2$, then:

$$\frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} \delta - 1}{(1 - c_1)e^\varepsilon + c_2 - 1} \leq \frac{1 - (1 - \delta) \frac{2}{e^\varepsilon + 1}}{1 - c_1}.$$

Because $c_1 < 1 < c_2$ and $e^\varepsilon \geq 1$, the denominator $(1 - c_1)e^\varepsilon + c_2 - 1$ is strictly positive, so we can multiply both sides of the inequality by it without changing the sign.

Set

$$A := (1 - c_1)e^\varepsilon + c_2 - 1 > 0.$$

The claim is equivalent to

$$(1 - c_1)(e^\varepsilon - 1) + (c_2 - c_1)\delta \leq \left(1 - (1 - \delta)\frac{2}{e^\varepsilon + 1}\right)A.$$

Bring the left-hand side to the right and factor out $(1 - \delta)$:

$$\begin{aligned} 0 &\leq A - (1 - \delta)\frac{2A}{e^\varepsilon + 1} - [(1 - c_1)(e^\varepsilon - 1) + (c_2 - c_1)\delta] \\ &= (c_2 - c_1)(1 - \delta) - (1 - \delta)\frac{2A}{e^\varepsilon + 1} \\ &= (1 - \delta)\left[(c_2 - c_1) - \frac{2A}{e^\varepsilon + 1}\right]. \end{aligned}$$

Since $1 - \delta \geq 0$, we only need to prove

$$(c_2 - c_1)(e^\varepsilon + 1) \geq 2A.$$

Substituting A gives

$$(c_2 - c_1)e^\varepsilon + (c_2 - c_1) \geq 2(1 - c_1)e^\varepsilon + 2(c_2 - 1).$$

Rearranging, we have:

$$(c_2 + c_1 - 2)(e^\varepsilon - 1) \geq 0.$$

Because $e^\varepsilon - 1 \geq 0$ for every $\varepsilon \geq 0$, the last inequality holds precisely when $c_1 + c_2 - 2 \geq 0$, i.e. when $c_1 + c_2 \geq 2$. This is exactly the hypothesis, so the desired inequality is proved and we have:

$$\lambda_{c_1, c_2, g_\varepsilon, \delta}^* = \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1}\delta - 1}{(1 - c_1)e^\varepsilon + c_2 - 1}.$$

Step 3: Optimality proof

Following Theorem 3.6, for the obtained value of $\lambda_{c_1, c_2, g_\varepsilon, \delta}$, the sampler

$$\mathbf{Q}_{c_1, c_2, \mu, g_\varepsilon, \delta}^*(P) := \lambda_{c_1, c_2, g_\varepsilon, \delta}^* P + (1 - \lambda_{c_1, c_2, g_\varepsilon, \delta}^*)\mu$$

belongs to the class $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g_\varepsilon, \delta}$ and is minimax-optimal with respect to any f -divergence D_f , that is,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, \mu, g_\varepsilon, \delta}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g_\varepsilon, \delta}, \tilde{\mathcal{P}}, f).$$

□

D Proofs of the main results

D.1 Proof of Proposition 3.3

Proof. Based on Proposition 6 in Dong et al. [21], for a trade-off function g , a sampler is g -FLDP if and only if it is $(\varepsilon, 1 + g^*(-e^\varepsilon))$ -LDP for all $\varepsilon \geq 0$. In the proof of Proposition C.2, we showed that if $\delta > \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1}$, then the identity sampler $\mathbf{Q}^I(P) = P$ satisfies (ε, δ) -LDP and has the trivial minimax risk zero.

Suppose the non-triviality condition (3) does not hold, i.e., for any $\varepsilon \geq 0$, we have:

$$1 + g^*(-e^\varepsilon) > \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1}.$$

In this case, the trivial sampler $\mathbf{Q}^I(P) = P$ satisfies $(\varepsilon, 1 + g^*(-e^\varepsilon))$ -LDP for any $\varepsilon \geq 0$ and therefore satisfies g -FLDP and has the trivial minimax risk zero. □

D.2 Proof of Theorem 3.6

Proof. In this proof, we assume that the non-triviality condition (3) is satisfied.

The proof of this theorem is directly based on Proposition C.2. Accordingly, the global minimax-optimal sampler in Theorem 3.6 builds upon the minimax-optimal sampler proposed in Proposition C.2. The lower bound in the optimality proof of Proposition C.2 relies on the existence of disjoint sets $A_1, \dots, A_t$, each with measure $\mu(A_i) = \alpha$, where $t = \frac{1}{\alpha}$. In fact, the converse part of the optimality proof hinges on the $(\alpha, \frac{1}{\alpha}, 1)$ -decomposability of μ , where $\frac{1}{\alpha} = \frac{c_2 - c_1}{1 - c_1} \in \mathbb{N}$. This technical requirement motivates the assumption $\frac{1}{\alpha} = \frac{c_2 - c_1}{1 - c_1} \in \mathbb{N}$ in Assumption 3.2.

Step 1: To propose a g -FLDP sampler $\mathbf{Q}_{c_1, c_2, \mu, g}^*$

From Proposition C.2, we know that for each pair (ε, δ) , the sampler

$$\mathbf{Q}_{\varepsilon, \delta}^*(P) = \lambda_{\varepsilon, \delta}^* P + (1 - \lambda_{\varepsilon, \delta}^*) \mu,$$

with

$$\lambda_{\varepsilon, \delta}^* = \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} \delta - 1}{(1 - c_1) e^\varepsilon + c_2 - 1},$$

is minimax-optimal in the class $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta}$.

Next, based on Proposition 6 in Dong et al. [21], for a trade-off function g , a sampler is g -FLDP if and only if it is $(\varepsilon, 1 + g^*(-e^\varepsilon))$ -LDP for all $\varepsilon \geq 0$. Consequently, if we set $\delta(\varepsilon) = 1 + g^*(-e^\varepsilon)$ and define

$$\mathbf{Q}_{\varepsilon, \delta(\varepsilon)}^*(P) := \lambda_{\varepsilon, \delta(\varepsilon)}^* P + (1 - \lambda_{\varepsilon, \delta(\varepsilon)}^*) \mu,$$

then by Proposition C.2, we have $\mathbf{Q}_{\varepsilon, \delta(\varepsilon)}^* \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta(\varepsilon)}$.

We then use Lemma C.1, which states that if $\mathbf{Q}_{\lambda_1}(P) = \lambda_1 P + (1 - \lambda_1) \mu$ is (ε, δ) -LDP for $0 \leq \lambda_1 \leq 1$, then for any $\lambda_2 \in [0, \lambda_1]$, the sampler $\mathbf{Q}_{\lambda_2}(P) = \lambda_2 P + (1 - \lambda_2) \mu$ is also (ε, δ) -LDP.

Therefore, if we define

$$\lambda_{c_1, c_2, g}^* = \inf_{\beta \geq 0} \frac{e^\beta + \frac{c_2 - c_1}{1 - c_1} [1 + g^*(-e^\beta)] - 1}{(1 - c_1) e^\beta + (c_2 - 1)},$$

this guarantees that the sampler

$$\mathbf{Q}_{c_1, c_2, \mu, g}^*(P) := \lambda_{c_1, c_2, g}^* P + (1 - \lambda_{c_1, c_2, g}^*) \mu$$

is $(\varepsilon, 1 + g^*(-e^\varepsilon))$ -LDP for all $\varepsilon \geq 0$. Note that the non-triviality constraint (3) guarantees that there exists an $\varepsilon \geq 0$ for which $1 + g^*(-e^\varepsilon) \leq \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1}$ and therefore $\lambda_{c_1, c_2, g}^* \leq 1$. By Proposition 6 of Dong et al. [21], being $(\varepsilon, 1 + g^*(-e^\varepsilon))$ -LDP for all $\varepsilon \geq 0$ is exactly the definition of being g -FLDP. Hence, $\mathbf{Q}_{c_1, c_2, \mu, g}^*$ is indeed a g -FLDP sampler and therefore belongs to $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}$. This completes the proof of the first step.

Step 2: To prove the optimality of $\mathbf{Q}_{c_1, c_2, \mu, g}^*$

Proposition 6 of Dong et al. [21] tells us that if we set

$$\delta(\varepsilon) = 1 + g^*(-e^\varepsilon),$$

then for every $\varepsilon \geq 0$ we have

$$\inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)) \geq \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon, \delta(\varepsilon)}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)). \quad (24)$$

On the other hand, define the sampler

$$\mathbf{Q}_{c_1, c_2, \mu, g}^*(P) := \lambda_{c_1, c_2, g}^* P + (1 - \lambda_{c_1, c_2, g}^*) \mu,$$

where

$$\lambda_{c_1, c_2, g}^* = \inf_{\varepsilon \geq 0} \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} (1 + g^*(-e^\varepsilon)) - 1}{(1 - c_1) e^\varepsilon + (c_2 - 1)}.$$

Let ε_g^* be the value of ε that achieves this infimum. The infimum is attained at a finite value of ε because, if we define

$$h_g(\varepsilon) := \frac{e^\varepsilon + \frac{c_2 - c_1}{1 - c_1} (1 + g^*(-e^\varepsilon)) - 1}{(1 - c_1)e^\varepsilon + (c_2 - 1)},$$

then we have

$$h_g(0) \leq \lim_{\varepsilon \rightarrow \infty} h_g(\varepsilon).$$

Together with the continuity of $h_g(\varepsilon)$, this implies that the infimum is achieved at some finite ε_g^* .

To establish the continuity of $h_g(\varepsilon)$, we note that both the numerator and denominator are continuous functions of ε , and the denominator is never zero. To verify that the numerator is continuous, it suffices to show that g^* is continuous. This follows directly from the definition of the convex conjugate and the fact that the original function g is continuous on the compact interval $[0, 1]$ and takes values in $[0, 1]$ (see Proposition B.2). Hence,

$$\lambda_{c_1, c_2, g}^* = \frac{e^{\varepsilon_g^*} + \frac{c_2 - c_1}{1 - c_1} (1 + g^*(-e^{\varepsilon_g^*})) - 1}{(1 - c_1)e^{\varepsilon_g^*} + (c_2 - 1)}.$$

By construction, $\mathbf{Q}_{c_1, c_2, \mu, g}^*$ belongs to the set $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon_g^*, \delta(\varepsilon_g^*)}$. From Proposition C.2, we know $\mathbf{Q}_{c_1, c_2, \mu, g}^*$ is minimax-optimal in $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon_g^*, \delta(\varepsilon_g^*)}$. Hence,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, \mu, g}^*(P)) = \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon_g^*, \delta(\varepsilon_g^*)}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)). \quad (25)$$

Finally, combining (24), (25), and the trivial inequality

$$\inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)) \leq \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, \mu, g}^*(P)),$$

we conclude:

$$\begin{aligned} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, \mu, g}^*(P)) &= \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon_g^*, \delta(\varepsilon_g^*)}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)) \\ &= \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon_g^*, \delta(\varepsilon_g^*)}, \tilde{\mathcal{P}}, f) \\ &= \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)) \\ &= \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f), \end{aligned}$$

which completes the proof of the second step.

Step 3: Computing the optimal value of $\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f)$

It follows from the previous step that in order to compute $\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f)$, it suffices to compute $\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon_g^*, \delta(\varepsilon_g^*)}, \tilde{\mathcal{P}}, f)$.

From Proposition C.2, we know that

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon_g^*, \delta(\varepsilon_g^*)}, \tilde{\mathcal{P}}, f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1)$$

for

$$r_1 = \frac{c_1}{c_2 - c_1} \cdot \frac{(1 - c_1)e^{\varepsilon_g^*} + c_2 - 1}{1 - \delta(\varepsilon_g^*)}, \quad r_2 = \frac{c_2}{c_2 - c_1} \cdot \frac{(1 - c_1)e^{\varepsilon_g^*} + c_2 - 1}{e^{\varepsilon_g^*} + \frac{c_2 - 1}{1 - c_1} \delta(\varepsilon_g^*)}.$$

We now compute r_1 and r_2 in terms of $\lambda_{c_1, c_2, g}^*$. For $\delta(\varepsilon_g^*) = 1 + g^*(-e^{\varepsilon_g^*})$, we have

$$r_1 = \frac{c_1}{c_2 - c_1} \cdot \frac{(1 - c_1)e^{\varepsilon_g^*} + c_2 - 1}{-g^*(-e^{\varepsilon_g^*})}, \quad r_2 = \frac{c_2}{c_2 - c_1} \cdot \frac{(1 - c_1)e^{\varepsilon_g^*} + c_2 - 1}{e^{\varepsilon_g^*} + \left(\frac{c_2 - 1}{1 - c_1}\right) (1 + g^*(-e^{\varepsilon_g^*}))}.$$

Moreover, we have

$$\lambda_{c_1, c_2, g}^* = \frac{e^{\varepsilon_g^*} + \frac{c_2 - c_1}{1 - c_1} \left(1 + g^*(-e^{\varepsilon_g^*})\right) - 1}{(1 - c_1)e^{\varepsilon_g^*} + (c_2 - 1)}.$$

Let $\theta := e^{\varepsilon_g^*}$, $\delta(\varepsilon_g^*) := 1 + g^*(-\theta)$, we have

$$\lambda_{c_1, c_2, g}^* = \frac{\theta + \frac{c_2 - c_1}{1 - c_1} \delta(\varepsilon_g^*) - 1}{(1 - c_1)\theta + (c_2 - 1)}. \quad (26)$$

Therefore,

$$r_1 = \frac{c_1}{c_2 - c_1} \frac{(1 - c_1)\theta + c_2 - 1}{1 - \delta(\varepsilon_g^*)}.$$

From (26),

$$1 - (1 - c_1)\lambda_{c_1, c_2, g}^* = \frac{(c_2 - c_1)[1 - \delta(\varepsilon_g^*)]}{(1 - c_1)\theta + (c_2 - 1)}.$$

Hence

$$r_1 = \frac{c_1}{1 - (1 - c_1)\lambda_{c_1, c_2, g}^*}.$$

Similarly,

$$r_2 = \frac{c_2}{c_2 - c_1} \frac{(1 - c_1)\theta + c_2 - 1}{\theta + \frac{c_2 - 1}{1 - c_1} \delta(\varepsilon_g^*)}.$$

Using (26) again,

$$(c_2 - 1)\lambda_{c_1, c_2, g}^* + 1 = \frac{c_2 - c_1}{(1 - c_1)\theta + (c_2 - 1)} \left[\theta + \frac{c_2 - 1}{1 - c_1} \delta(\varepsilon_g^*) \right].$$

Thus

$$r_2 = \frac{c_2}{(c_2 - 1)\lambda_{c_1, c_2, g}^* + 1}.$$

In conclusion,

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon_g^*, \delta(\varepsilon_g^*)}, \tilde{\mathcal{P}}, f) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1),$$

for

$$r_1 = \frac{c_1}{1 - (1 - c_1)\lambda_{c_1, c_2, g}^*} \quad \text{and} \quad r_2 = \frac{c_2}{(c_2 - 1)\lambda_{c_1, c_2, g}^* + 1}.$$

□

D.3 Proof of Theorem 3.4

Proof. Based on Theorem 3.6, we define $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$ under Assumption 3.2. Moreover, suppose μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1 - c_1}{c_2 - c_1}$. Then, the sampler defined as

$$\mathbf{Q}_{c_1, c_2, \mu, g}^*(P) = \lambda_{c_1, c_2, g}^* P + (1 - \lambda_{c_1, c_2, g}^*) \mu$$

satisfies g -FLDP and is minimax-optimal with respect to any f -divergence, where $\lambda_{c_1, c_2, g}^*$ is defined as

$$\lambda_{c_1, c_2, g}^* = \inf_{\beta \geq 0} \frac{e^\beta + \frac{c_2 - c_1}{1 - c_1} (1 + g^*(-e^\beta)) - 1}{(1 - c_1)e^\beta + c_2 - 1}.$$

Theorem 3.4 is a special case of Theorem 3.6. Below, we demonstrate this reduction explicitly.

Recall that in the setup of Theorem 3.6, for a general sample space $\mathcal{X}$, the universe $\tilde{\mathcal{P}}$ is defined as

$$\tilde{\mathcal{P}}_{c_1, c_2, \mu} := \left\{ P \in \mathcal{P}(\mathcal{X}) : P \ll \mu, c_1 \leq \frac{dP}{d\mu} \leq c_2, \mu\text{-a.e.} \right\}.$$

In the setup of Theorem 3.4, $\mathcal{X} = \mathbb{R}^n$, and $\tilde{\mathcal{P}}$ is defined as

$$\tilde{\mathcal{P}}_{c_1, c_2, h} := \{P \in \mathcal{C}(\mathbb{R}^n) : c_1 h(x) \leq p(x) \leq c_2 h(x), \quad \forall x \in \mathbb{R}^n\}.$$

Let λ denote the Lebesgue measure on $\mathbb{R}^n$. We have $\mathcal{X} = \mathbb{R}^n$, and $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, h} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$, where $\mu \ll \lambda$ and $\frac{d\mu}{d\lambda} = h$. Therefore, $\mathbf{Q}_{c_1, c_2, \mu, g}^* = \mathbf{Q}_{c_1, c_2, h, g}^*$.

This is because for each $P \in \tilde{\mathcal{P}}$ with corresponding PDF $p(x)$, the chain rule of the Radon-Nikodym derivative gives:

$$p(x) = \frac{dP}{d\lambda} = \frac{dP}{d\mu}(x) \cdot \frac{d\mu}{d\lambda}(x) = \frac{dP}{d\mu}(x)h(x).$$

It remains to show that μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable for $\alpha = \frac{1-c_1}{c_2-c_1}$. From Assumption 3.2, we know that $t = \frac{c_2-c_1}{1-c_1} \in \mathbb{N}$. Therefore we need to show that μ is $(\frac{1}{t}, t, 1)$ -decomposable.

Since $\mu \ll \lambda$, the function $s \mapsto \mu((-\infty, s] \times \mathbb{R}^{n-1})$ is continuous. As $s \rightarrow -\infty$, this measure tends to 0, and as $s \rightarrow \infty$, it tends to 1. By the intermediate value theorem, for each $i \in [t]$, there exists a threshold $s_i \in \mathbb{R}$ such that

$$\mu((-\infty, s_i] \times \mathbb{R}^{n-1}) = \alpha i.$$

We then define the sets $A_1 = (-\infty, s_1] \times \mathbb{R}^{n-1}$, and for $i \geq 2$, set $A_i = (s_{i-1}, s_i] \times \mathbb{R}^{n-1}$. These sets satisfy the requirements of $(\frac{1}{t}, t, 1)$ -decomposability. Therefore, μ is indeed $(\frac{1}{t}, t, 1)$ -decomposable, completing the proof.

Therefore, Theorem 3.6 contains Theorem 3.4 as a special case. Under Assumption 3.2, the sampler $\mathbf{Q}_{c_1, c_2, h, g}^*$, defined as a continuous distribution whose density is given by

$$q_g^*(P)(x) := \lambda_{c_1, c_2, g}^* p(x) + (1 - \lambda_{c_1, c_2, g}^*) h(x), \quad \lambda_{c_1, c_2, g}^* = \inf_{\beta \geq 0} \frac{e^\beta + \frac{c_2-c_1}{1-c_1}(1+g^*(-e^\beta))-1}{(1-c_1)e^\beta + c_2 - 1}, \quad (27)$$

satisfies g -FLDP and is minimax-optimal under any f -divergence. That is,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, h, g}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f) = \frac{1-r_1}{r_2-r_1} f(r_2) + \frac{r_2-1}{r_2-r_1} f(r_1),$$

for $r_1 = \frac{c_1}{1-(1-c_1)\lambda_{c_1, c_2, g}^*}$ and $r_2 = \frac{c_2}{(c_2-1)\lambda_{c_1, c_2, g}^*+1}$. $\square$

D.4 Proof of Theorem 3.5

Proof. From Theorem 3.6, we know that if $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$ is defined under Assumption 3.2 and μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1-c_1}{c_2-c_1}$, then the sampler $\mathbf{Q}_{c_1, c_2, \mu, g}^*(P) := \lambda_{c_1, c_2, g}^* P + (1 - \lambda_{c_1, c_2, g}^*) \mu$ belongs to the class $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}$ and is minimax-optimal under any f -divergence D_f ; that is,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, \mu, g}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f),$$

In this case,

$$\lambda_{c_1, c_2, g}^* = \inf_{\beta \geq 0} \frac{e^\beta + \frac{c_2-c_1}{1-c_1}(1+g^*(-e^\beta))-1}{(1-c_1)e^\beta + c_2 - 1}.$$

It can be shown that:

$$\mathcal{P}([k]) = \tilde{\mathcal{P}}_{0, k, \mu_k} := \left\{ P \in \mathcal{P}([k]) : P \ll \mu_k, \quad 0 \leq \frac{dP}{d\mu_k} \leq k \quad \mu_k\text{-a.e.} \right\}$$

Consider the following setting:

$$\tilde{\mathcal{P}} = \mathcal{P}([k]), \quad \mathcal{X} = [k], \quad c_1 = 0, \quad c_2 = k, \quad \mu = \mu_k,$$

where μ_k is the uniform distribution on $[k]$. In this case, we have $\frac{c_2-c_1}{1-c_1} = k \in \mathbb{N}$, and μ_k is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1}{k}$. Moreover, for this specific instantiation, it can be shown that for all $\varepsilon \geq 0$,

$$1 + g^*(-e^\varepsilon) \leq \frac{(c_2 - c_1 e^\varepsilon)(1 - c_1)}{c_2 - c_1} = \frac{(k)(1)}{k} = 1.$$

This comes from the definition of the convex conjugate and the properties of the function g :

$$\forall \varepsilon \geq 0 : \quad -1 \leq g^*(-e^\varepsilon) \leq 0,$$

and hence,

$$1 + g^*(-e^\varepsilon) \leq 1.$$

Therefore, all conditions of Theorem 3.6 are satisfied and Theorem 3.5 is a special case of the more general case Theorem 3.6. Hence, the sampler $\mathbf{Q}_{k,g}^*(P) := \lambda_{k,g}^* P + (1 - \lambda_{k,g}^*) \mu_k$ belongs to the class $\mathcal{Q}_{[k], \tilde{\mathcal{P}}, g}$ and is minimax-optimal under any f -divergence D_f ; that is,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{k,g}^*(P)) = \mathcal{R}(\mathcal{Q}_{[k], \tilde{\mathcal{P}}, g}, \tilde{\mathcal{P}}, f),$$

where $\lambda_{k,g}^*$ is the optimal solution to the following optimization problem:

$$\lambda_{k,g}^* = \inf_{\beta \geq 0} \frac{e^\beta + k(1 + g^*(-e^\beta)) - 1}{e^\beta + k - 1}.$$

□

D.5 Proof of Corollary 3.7

Proof. Corollary 3.7 follows immediately from Corollary C.4 by setting $\delta = 0$. With this choice, we have $g_\varepsilon = g_{\varepsilon,0}$, thus the proof is identical. Moreover, as in the proof of Theorem 3.4, the continuous case of Corollary 3.7 can be shown to be a special instance of the more general result stated in Corollary C.4. For completeness, a brief proof sketch is provided below.

Step 1: Find g_ε^*

$$g_\varepsilon^*(y) = \begin{cases} -1, & \text{if } y < -e^\varepsilon, \\ \frac{y-1}{e^\varepsilon+1}, & \text{if } -e^\varepsilon \leq y \leq -e^{-\varepsilon}, \\ y, & \text{if } y > -e^{-\varepsilon}. \end{cases} \quad (28)$$

Step 2: Solve the infimum to obtain $\lambda_{c_1, c_2, g_\varepsilon}^*$

$$\lambda_{c_1, c_2, g_\varepsilon}^* = \begin{cases} \frac{e^\varepsilon - 1}{(e^\varepsilon + 1)(1 - c_1)}, & \text{if } c_1 + c_2 < 2, \\ \frac{e^\varepsilon - 1}{(1 - c_1)e^\varepsilon + c_2 - 1}, & \text{if } c_1 + c_2 \geq 2. \end{cases} \quad (29)$$

Since $\frac{c_2 - c_1}{1 - c_1} \in \mathbb{N}$ and $c_2 > 1$, it follows that $\frac{c_2 - c_1}{1 - c_1} \geq 2$, which in turn implies $c_1 + c_2 \geq 2$. Hence:

$$\lambda_{c_1, c_2, g_\varepsilon}^* = \frac{e^\varepsilon - 1}{(1 - c_1)e^\varepsilon + c_2 - 1}.$$

Step 3: Optimality proof

Following Theorem 3.6, for the obtained value of $\lambda_{c_1, c_2, g_\varepsilon}^*$, the sampler

$$\mathbf{Q}_{c_1, c_2, \mu, g_\varepsilon}^*(P) := \lambda_{c_1, c_2, g_\varepsilon}^* P + (1 - \lambda_{c_1, c_2, g_\varepsilon}^*) \mu$$

belongs to the class $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g_\varepsilon, \delta}$ and is minimax-optimal with respect to any f -divergence D_f , that is,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, \mu, g_\varepsilon}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g_\varepsilon}, \tilde{\mathcal{P}}, f).$$

Therefore, we have:

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{c_1, c_2, h, g_\varepsilon}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, g_\varepsilon}, \tilde{\mathcal{P}}, f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1),$$

for $r_1 = c_1 \cdot \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{c_2 - c_1}$, and $r_2 = \frac{c_2}{c_2 - c_1} \cdot \frac{(1 - c_1)e^\varepsilon + c_2 - 1}{e^\varepsilon}$.

The optimal value of $\mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, g_\varepsilon}, \tilde{\mathcal{P}}, f)$ is obtained directly from Theorem 3.4 by substituting

$$\lambda_{c_1, c_2, g_\varepsilon}^* = \frac{e^\varepsilon - 1}{(1 - c_1)e^\varepsilon + c_2 - 1}.$$

□

D.6 Proof of Corollary 3.8

Proof. From Theorem 3.6 we know that for

$$\lambda_{G_\nu}^* = \inf_{\beta \geq 0} \frac{e^\beta + \frac{c_2 - c_1}{1 - c_1} (1 + G_\nu^*(-e^\beta)) - 1}{(1 - c_1)e^\beta + c_2 - 1}.$$

the sampler (4) for $g = G_\nu$ belongs to $\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, G_\nu}$ and is minimax-optimal with respect to any f -divergence, that is,

$$\sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}_{G_\nu}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathbb{R}^n, \tilde{\mathcal{P}}, G_\nu}, \tilde{\mathcal{P}}, f).$$

Now, we compute $G_\nu^*(-e^\beta)$. It is shown in Corollary 1 of Dong et al. [21] that

$$G_\nu^*(y) = y \Phi\left(-\frac{\nu}{2} - \frac{1}{\nu} \log(-y)\right) - \Phi\left(-\frac{\nu}{2} + \frac{1}{\nu} \log(-y)\right).$$

When $y = -e^\beta$,

$$G_\nu^*(-e^\beta) = -e^\beta \Phi\left(-\frac{\nu}{2} - \frac{\beta}{\nu}\right) - \Phi\left(-\frac{\nu}{2} + \frac{\beta}{\nu}\right)$$

Therefore, we have:

$$\begin{aligned} \lambda_{G_\nu}^* &= \inf_{\beta \geq 0} \frac{e^\beta + \left(\frac{c_2 - c_1}{1 - c_1}\right) \left(1 - e^\beta \Phi\left(-\frac{\nu}{2} - \frac{\beta}{\nu}\right) - \Phi\left(-\frac{\nu}{2} + \frac{\beta}{\nu}\right)\right) - 1}{(1 - c_1)e^\beta + c_2 - 1} \\ &= \inf_{\beta \geq 0} \frac{e^\beta + \left(\frac{c_2 - c_1}{1 - c_1}\right) \left(\Phi\left(\frac{\nu}{2} - \frac{\beta}{\nu}\right) - e^\beta \Phi\left(-\frac{\nu}{2} - \frac{\beta}{\nu}\right)\right) - 1}{(1 - c_1)e^\beta + c_2 - 1}. \end{aligned}$$

□

D.7 Proof of Theorem 4.1

Proof. Step 1: Lower bound

In the first step of the proof, we establish a lower bound for the local minimax objective function $\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, N_\gamma(P_0), f)$ as follows:

$$\begin{aligned} \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, N_\gamma(P_0), f) &= \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}} \sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}(P)) \\ &\geq \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), g}} \sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}(P)) \\ &= \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), g}, N_\gamma(P_0), f). \end{aligned} \quad (30)$$

The inequality in the second line follows from the fact that any sampler $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$ in $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}$ also belongs to $\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), g}$. The third line is by definition of the global minimax risk in (1) when the universe $\tilde{\mathcal{P}}$ is restricted to $N_\gamma(P_0)$.

Step 2: Upper bound

We now construct a sampler that achieves the lower bound. Specifically, we aim to find a sampler $\mathbf{Q}_{g, N_\gamma(P_0)}^* \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}$ such that

$$\sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_{g, N_\gamma(P_0)}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), g}, N_\gamma(P_0), f).$$

Observe that the neighborhood $N_\gamma(P_0)$ aligns with the class $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$ defined in Theorem 3.6:

$$\tilde{\mathcal{P}}_{c_1, c_2, \mu} := \left\{ P \in \mathcal{P}(\mathcal{X}) : P \ll \mu, \quad c_1 \leq \frac{dP}{d\mu} \leq c_2 \quad \mu\text{-a.e.} \right\},$$

$$N_\gamma(P_0) := \left\{ P \in \mathcal{P}(\mathcal{X}) : \frac{1}{\gamma} \leq \frac{dP}{dP_0}(x) \leq \gamma \quad \forall x \in \mathcal{X} \right\}.$$

By substituting $c_1 = \frac{1}{\gamma}$, $c_2 = \gamma$, and $\mu = P_0$, we obtain $\tilde{\mathcal{P}}_{c_1, c_2, \mu} = N_\gamma(P_0)$.

Additionally, since $\gamma \in \mathbb{N}$, it follows that $\frac{c_2 - c_1}{1 - c_1} = \frac{\gamma - \frac{1}{\gamma}}{1 - \frac{1}{\gamma}} \in \mathbb{N}$, satisfying the integer condition in Theorem 3.6. The $(\alpha, \frac{1}{\alpha}, 1)$ -decomposability of P_0 with $\alpha = \frac{1}{\gamma+1}$ also matches the requirement that μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1 - c_1}{c_2 - c_1}$.

Therefore, all the conditions in Theorem 3.6 are satisfied. Thus, the optimal sampler $\mathbf{Q}_g^*$ corresponding to $(c_1, c_2, \mu) = (\frac{1}{\gamma}, \gamma, P_0)$ achieves

$$\sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_g^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), g}, N_\gamma(P_0), f).$$

Since $\mathbf{Q}_g^*$ is defined only on $N_\gamma(P_0)$, we extend it to the larger domain $\tilde{\mathcal{P}}$ by defining:

$$\mathbf{Q}_{g, N_\gamma(P_0)}^*(P) := \begin{cases} \mathbf{Q}_g^*(P), & \text{if } P \in N_\gamma(P_0), \\ \mathbf{Q}_g^*(\hat{P}), & \text{otherwise,} \end{cases}$$

where $\hat{P} \in N_\gamma(P_0)$ is a distribution that minimizes $D_f(P \| P')$ over all $P' \in N_\gamma(P_0)$.

By construction, $\mathbf{Q}_{g, N_\gamma(P_0)}^* \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}$, and for all $P \in N_\gamma(P_0)$, we have:

$$\begin{aligned} \sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_{g, N_\gamma(P_0)}^*(P)) &= \sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_g^*(P)) \\ &= \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), g}, N_\gamma(P_0), f). \end{aligned}$$

Hence, we obtain the upper bound:

$$\begin{aligned} \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, N_\gamma(P_0), f) &= \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}} \sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}(P)) \\ &\leq \sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_{g, N_\gamma(P_0)}^*(P)) \\ &= \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), g}, N_\gamma(P_0), f). \end{aligned} \quad (31)$$

Combining equations (30) and (31), we conclude that

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, g}, N_\gamma(P_0), f) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), g}, N_\gamma(P_0), f).$$

Note that this proof is stated for a general space $\mathcal{X}$ and a general universe $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$. It extends directly to the continuous case with $\mathcal{X} = \mathbb{R}^n$ and $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, h}$, following the same reduction argument used in the proof of Theorem 3.4. $\square$

D.8 Proof of Theorem 5.1

Proof. let the global minimax formulation under ε -LDP be defined as:

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, \tilde{\mathcal{P}}, f) := \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}} \sup_{P \in \tilde{\mathcal{P}}} D_f(P \| \mathbf{Q}(P)) \quad (32)$$

Under certain assumptions, formally defined in Section 3, Park et al. [17] derives the optimal sampler and optimal value for the optimization problem (32).

Step 1: Lower bound

In the first step of the proof, we establish a lower bound for the local minimax objective function $\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, N_\gamma(P_0), f)$ as follows:

$$\begin{aligned}
\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, N_\gamma(P_0), f) &= \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}} \sup_{P \in N_\gamma(P_0)} D_f(P \parallel \mathbf{Q}(P)) \\
&\geq \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}} \sup_{P \in N_\gamma(P_0)} D_f(P \parallel \mathbf{Q}(P)) \\
&= \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f).
\end{aligned} \tag{33}$$

The inequality in the second line follows from the fact that any sampler $\mathbf{Q} : \tilde{\mathcal{P}} \rightarrow \mathcal{P}(\mathcal{X})$ in $\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}$ also belongs to $\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}$. The third line is by definition of the global minimax risk in (32) when the universe $\tilde{\mathcal{P}}$ is restricted to $N_\gamma(P_0)$.

Step 2: Upper bound

We now construct a sampler that achieves the lower bound. Specifically, we aim to find a sampler $\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^* \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}$ such that

$$\sup_{P \in N_\gamma(P_0)} D_f(P \parallel \mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f).$$

Let sampler $\mathbf{Q}_{\varepsilon}^* : N_\gamma(P_0) \rightarrow \mathcal{P}(\mathcal{X})$ be an optimal sampler for the global minimax problem (32) where the universe $\tilde{\mathcal{P}}$ is restricted to $N_\gamma(P_0)$. In other words, $\mathbf{Q}_{\varepsilon}^*$ achieves

$$\sup_{P \in N_\gamma(P_0)} D_f(P \parallel \mathbf{Q}_{\varepsilon}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f).$$

Since $\mathbf{Q}_{\varepsilon}^*$ is defined only on $N_\gamma(P_0)$, we extend it to the larger domain $\tilde{\mathcal{P}}$ by defining:

$$\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*(P) := \begin{cases} \mathbf{Q}_{\varepsilon}^*(P), & \text{if } P \in N_\gamma(P_0), \\ \mathbf{Q}_{\varepsilon}^*(\hat{P}), & \text{otherwise,} \end{cases} \tag{34}$$

where $\hat{P} \in N_\gamma(P_0)$ is a distribution that minimizes $D_f(P \parallel P')$ over all $P' \in N_\gamma(P_0)$.

By construction, $\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^* \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}$, and for all $P \in N_\gamma(P_0)$, we have:

$$\begin{aligned}
\sup_{P \in N_\gamma(P_0)} D_f(P \parallel \mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*(P)) &= \sup_{P \in N_\gamma(P_0)} D_f(P \parallel \mathbf{Q}_{\varepsilon}^*(P)) \\
&= \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f).
\end{aligned}$$

Hence, we obtain the upper bound:

$$\begin{aligned}
\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, N_\gamma(P_0), f) &= \inf_{\mathbf{Q} \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}} \sup_{P \in N_\gamma(P_0)} D_f(P \parallel \mathbf{Q}(P)) \\
&\leq \sup_{P \in N_\gamma(P_0)} D_f(P \parallel \mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*(P)) \\
&= \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f).
\end{aligned} \tag{35}$$

Combining equations (33) and (35), we conclude that

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, N_\gamma(P_0), f) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f).$$

Step 3: Local minimax risk in a closed form

Under the same non-triviality assumption as Park et al. [17], for the pure LDP case we assume $\varepsilon < 2 \log(\gamma)$; otherwise, the identity sampler becomes the trivial local minimax-optimal sampler with zero minimax risk.

Consider the case where $2 \log(\gamma) \leq \varepsilon$. In this case, for any $P_1, P_2 \in N_\gamma(P_0)$, we have $p_1(x)/p_2(x) \leq \frac{\gamma}{1} < e^\varepsilon$, hence we can easily observe that the sampler $\mathbf{Q}_\varepsilon^I$ defined as $\mathbf{Q}_\varepsilon^I(P) = P$

for all $P \in N_\gamma(P_0)$ satisfies ε -LDP and $\mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f) = 0$. Since $\mathbf{Q}_\varepsilon^I$ is defined only on $N_\gamma(P_0)$, we extend it to the larger domain $\tilde{\mathcal{P}}$ by defining:

$$\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^I(P) := \begin{cases} \mathbf{Q}_\varepsilon^I(P), & \text{if } P \in N_\gamma(P_0), \\ \mathbf{Q}_\varepsilon^I(\hat{P}), & \text{otherwise,} \end{cases}$$

where $\hat{P} \in N_\gamma(P_0)$ is a distribution that minimizes $D_f(P \| P')$ over all $P' \in N_\gamma(P_0)$.

We can conclude $\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, N_\gamma(P_0), f) = 0$. Therefore, to exclude trivial samplers we assume $\varepsilon < 2 \log(\gamma)$.

In order to obtain the local minimax risk under ε -LDP, we refer to Theorem C.4 in Park et al. [17].

Theorem D.1 (Theorem C.4 of Park et al. [17]). *Let $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$ be a set of distributions such that the normalization condition (14) holds, and suppose μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1-c_1}{c_2-c_1}$. Define*

$$b := \frac{c_2 - c_1}{(e^\varepsilon - 1)(1 - c_1) + c_2 - c_1}, \quad r_1 := \frac{c_1}{b}, \quad r_2 := \frac{c_2}{be^\varepsilon}.$$

Then, the optimal value of the problem (32) is given by

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, \tilde{\mathcal{P}}, f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1).$$

Furthermore, the sampler $\mathbf{Q}_{c_1, c_2, \mu, \varepsilon}^$ defined below is an optimal solution to problem (32) under any f -divergence D_f . For each $P \in \tilde{\mathcal{P}}$, let $\mathbf{Q}_{c_1, c_2, \mu, \varepsilon}^*(P) = Q$ be a probability measure $Q \ll \mu$, such that*

$$\frac{dQ}{d\mu}(x) = \text{clip} \left(\frac{1}{r_P} \frac{dP}{d\mu}(x) ; b, be^\varepsilon \right),$$

where r_P is the normalizing constant ensuring that $\int \frac{dQ}{d\mu} d\mu(x) = 1$.

Observe that the neighborhood $N_\gamma(P_0)$ aligns with the class $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$ defined in Theorem 3.6:

$$\begin{aligned} \tilde{\mathcal{P}}_{c_1, c_2, \mu} &:= \left\{ P \in \mathcal{P}(\mathcal{X}) : P \ll \mu, \quad c_1 \leq \frac{dP}{d\mu} \leq c_2 \quad \mu\text{-a.e.} \right\}, \\ N_\gamma(P_0) &:= \left\{ P \in \mathcal{P}(\mathcal{X}) : \frac{1}{\gamma} \leq \frac{dP}{dP_0}(x) \leq \gamma \quad \forall x \in \mathcal{X} \right\}. \end{aligned}$$

By substituting $c_1 = \frac{1}{\gamma}$, $c_2 = \gamma$, and $\mu = P_0$, we obtain $\tilde{\mathcal{P}}_{c_1, c_2, \mu} = N_\gamma(P_0)$.

Additionally, the $(\alpha, \frac{1}{\alpha}, 1)$ -decomposability of P_0 with $\alpha = \frac{1}{\gamma+1}$ also matches the requirement that μ is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable with $\alpha = \frac{1-c_1}{c_2-c_1}$. We now aim to compute the closed-form expression for

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f).$$

By the equivalence $\tilde{\mathcal{P}}_{c_1, c_2, \mu} = N_\gamma(P_0)$ and Theorem D.1 applied with $(c_1, c_2, \mu) = (\frac{1}{\gamma}, \gamma, P_0)$, we have:

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1),$$

where

$$b := \frac{\gamma - \frac{1}{\gamma}}{(e^\varepsilon - 1)(1 - \frac{1}{\gamma}) + \gamma - \frac{1}{\gamma}} = \frac{\gamma + 1}{\gamma + e^\varepsilon}, \quad r_1 := \frac{1}{\gamma b} = \frac{e^\varepsilon + \gamma}{\gamma(\gamma + 1)}, \quad r_2 := \frac{\gamma}{be^\varepsilon} = \frac{\gamma(e^\varepsilon + \gamma)}{e^\varepsilon(\gamma + 1)}.$$

Since

$$\mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, N_\gamma(P_0), f) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f),$$

the optimal minimax risk is obtained.

Step 4: Optimal sampler

From the upper bound proof of the theorem, we know that the sampler $\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*$, defined in (34), is an optimal sampler, provided that

$$\sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_\varepsilon^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}, N_\gamma(P_0), f).$$

Therefore, it suffices to construct a sampler $\mathbf{Q}_\varepsilon^* : N_\gamma(P_0) \rightarrow \mathcal{P}(\mathcal{X})$ that belongs to $\mathcal{Q}_{\mathcal{X}, N_\gamma(P_0), \varepsilon}$ and satisfies the above equality.

By applying Theorem D.1 with $(c_1, c_2, \mu) = (\frac{1}{\gamma}, \gamma, P_0)$, we obtain the following definition for such a sampler. For each $P \in N_\gamma(P_0)$, define $\mathbf{Q}_\varepsilon^*(P) = Q$, where Q is a probability measure satisfying $Q \ll P_0$ and given by

$$\frac{dQ}{dP_0}(x) = \text{clip} \left(\frac{1}{r_P} \frac{dP}{dP_0}(x); \frac{\gamma+1}{\gamma+e^\varepsilon}, \frac{\gamma+1}{\gamma+e^\varepsilon} e^\varepsilon \right),$$

with r_P denoting the normalizing constant ensuring that $\int \frac{dQ}{dP_0} dP_0(x) = 1$.

We extend this sampler to a function over all of $\mathcal{P}(\mathcal{X})$ by defining

$$\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*(P) := \begin{cases} \mathbf{Q}_\varepsilon^*(P), & \text{if } P \in N_\gamma(P_0), \\ \mathbf{Q}_\varepsilon^*(\hat{P}), & \text{otherwise,} \end{cases}$$

where $\hat{P} \in N_\gamma(P_0)$ is a distribution that minimizes $D_f(P \| P')$ over all $P' \in N_\gamma(P_0)$.

Then $\mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^* \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}$, and the following holds:

$$\sup_{P \in N_\gamma(P_0)} D_f(P \| \mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*(P)) = \mathcal{R}(\mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}, N_\gamma(P_0), f).$$

Note that this proof is stated for a general space $\mathcal{X}$ and a general universe $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, \mu}$. It extends directly to the continuous case with $\mathcal{X} = \mathbb{R}^n$ and $\tilde{\mathcal{P}} = \tilde{\mathcal{P}}_{c_1, c_2, h}$, following the same reduction argument used in the proof of Theorem 3.4. $\square$

D.9 Proof of Proposition 5.2

Proof. To prove this result, we need to define preliminary samplers and the universe set. Let

$$\tilde{\mathcal{P}}_{c_1, c_2, \mu} := \left\{ P \in \mathcal{P}(\mathcal{X}) : P \ll \mu, \quad c_1 \leq \frac{dP}{d\mu} \leq c_2 \quad \mu\text{-a.e.} \right\},$$

Define a sampler $\mathbf{Q}_{c_1, c_2, \mu, \varepsilon}^* \in \mathcal{Q}_{\mathcal{X}, \tilde{\mathcal{P}}, \varepsilon}$ as follows:

For each $P \in \tilde{\mathcal{P}}$, $\mathbf{Q}_{c_1, c_2, \mu, \varepsilon}^*(P) := Q$ is defined as a probability measure such that $Q \ll \mu$ and

$$\frac{dQ}{d\mu}(x) = \text{clip} \left(\frac{1}{r_P} \frac{dP}{d\mu}(x); b, be^\varepsilon \right),$$

where $r_P > 0$ is a constant depending on P so that $\int \frac{dQ}{d\mu} d\mu(x) = 1$.

Proposition D.2 (Proposition C.8. of Park et al. [17]). *For any $P \in \tilde{\mathcal{P}}_{c_1, c_2, \mu}$ and any f -divergences D_f , we have*

$$D_f(P \| \mathbf{Q}_{c_1, c_2, \mu, \varepsilon}^*(P)) = \inf_{\mathbf{Q} \in \tilde{\mathcal{P}}_{b, be^\varepsilon, \mu}} D_f(P \| \mathbf{Q}),$$

for

$$b = \frac{c_2 - c_1}{(e^\varepsilon - 1)(1 - c_1) + c_2 - c_1}.$$

We want to prove that if we define the neighborhood as

$$N_\gamma(P_0) := \{P \in \mathcal{P}(\mathcal{X}) : E_\gamma(P \| P_0) = E_\gamma(P_0 \| P) = 0\},$$

and let $\mathbf{Q}_\varepsilon^*$ denote the optimal sampler from Theorem 5.1, and let $\mathbf{Q}_{g_\varepsilon}^*$ be the instantiation of Theorem 4.1 with $g = g_\varepsilon$. Then, for all $P \in N_\gamma(P_0)$,

$$D_f(P \| \mathbf{Q}_\varepsilon^*) \leq D_f(P \| \mathbf{Q}_{g_\varepsilon}^*).$$

Note that this inequality is equivalent to the condition that for all $P \in N_\gamma(P_0)$,

$$D_f(P \| \mathbf{Q}_{\varepsilon, N_\gamma(P_0)}^*) \leq D_f(P \| \mathbf{Q}_{g_\varepsilon, N_\gamma(P_0)}^*),$$

which is the ultimate result we want to prove in this proposition.

Therefore, it suffices to show that for $\tilde{\mathcal{P}}_{c_1, c_2, \mu} = N_\gamma(P_0)$, we have $\mathbf{Q}_\varepsilon^* = \mathbf{Q}_{c_1, c_2, \mu, \varepsilon}^*$ and that $\mathbf{Q}_{g_\varepsilon}^* \in \tilde{\mathcal{P}}_{b, be^\varepsilon, \mu}$.

It follows that for $c_1 = \frac{1}{\gamma}$, $c_2 = \gamma$, and $\mu = P_0$, we have the equivalence $\tilde{\mathcal{P}}_{c_1, c_2, \mu} = N_\gamma(P_0)$. For this instantiation of $\tilde{\mathcal{P}}_{c_1, c_2, \mu}$, the optimal sampler $\mathbf{Q}_{c_1, c_2, \mu, \varepsilon}^*$ for input P is defined as a probability measure such that $Q \ll P_0$ and

$$\frac{dQ}{dP_0}(x) = \text{clip} \left(\frac{1}{r_P} \frac{dP}{dP_0}(x); b, be^\varepsilon \right),$$

where

$$b = \frac{\gamma - \frac{1}{\gamma}}{(e^\varepsilon - 1)(1 - \frac{1}{\gamma}) + \gamma - \frac{1}{\gamma}} = \frac{\gamma + 1}{e^\varepsilon + \gamma}.$$

This is exactly $\mathbf{Q}_\varepsilon^*$, the optimal sampler from Theorem 5.1. Thus, it remains to prove that $\mathbf{Q}_{g_\varepsilon}^* \in \tilde{\mathcal{P}}_{b, be^\varepsilon, P_0}$ for $b = \frac{\gamma+1}{e^\varepsilon+\gamma}$.

Combining Theorem 4.1 and Corollary 3.7, we obtain

$$\mathbf{Q}_{g_\varepsilon}^*(P) := \lambda_{g_\varepsilon}^* P + (1 - \lambda_{g_\varepsilon}^*) P_0, \quad \lambda_{g_\varepsilon}^* = \frac{e^\varepsilon - 1}{(1 - \frac{1}{\gamma})e^\varepsilon + \gamma - 1} = \frac{\gamma(e^\varepsilon - 1)}{(\gamma - 1)(e^\varepsilon + \gamma)}.$$

Recall that

$$\tilde{\mathcal{P}}_{b, be^\varepsilon, P_0} := \left\{ Q \in \mathcal{P}(\mathcal{X}) : Q \ll P_0, \quad b \leq \frac{dQ}{dP_0} \leq be^\varepsilon \quad P_0\text{-a.e.} \right\},$$

and

$$N_\gamma(P_0) := \left\{ P \in \mathcal{P}(\mathcal{X}) : P \ll P_0, \quad \frac{1}{\gamma} \leq \frac{dP}{dP_0} \leq \gamma \quad P_0\text{-a.e.} \right\}.$$

Therefore, if $P \in N_\gamma(P_0)$, we have

$$\frac{d\mathbf{Q}_{g_\varepsilon}^*}{dP_0} = \lambda_{g_\varepsilon}^* \frac{dP}{dP_0} + (1 - \lambda_{g_\varepsilon}^*).$$

It follows that $\mathbf{Q}_{g_\varepsilon}^* \ll P_0$ and

$$\frac{\lambda_{g_\varepsilon}^*}{\gamma} + (1 - \lambda_{g_\varepsilon}^*) \leq \frac{d\mathbf{Q}_{g_\varepsilon}^*}{dP_0} \leq \gamma \lambda_{g_\varepsilon}^* + (1 - \lambda_{g_\varepsilon}^*) \quad P_0\text{-a.e..}$$

Equivalently,

$$\frac{e^\varepsilon - 1}{(\gamma - 1)(e^\varepsilon + \gamma)} + \frac{\gamma^2 - e^\varepsilon}{(\gamma - 1)(e^\varepsilon + \gamma)} \leq \frac{d\mathbf{Q}_{g_\varepsilon}^*}{dP_0} \leq \frac{\gamma^2(e^\varepsilon - 1)}{(\gamma - 1)(e^\varepsilon + \gamma)} + \frac{\gamma^2 - e^\varepsilon}{(\gamma - 1)(e^\varepsilon + \gamma)} \quad P_0\text{-a.e..}$$

As a result,

$$\frac{\gamma^2 - 1}{(\gamma - 1)(e^\varepsilon + \gamma)} \leq \frac{d\mathbf{Q}_{g_\varepsilon}^*}{dP_0} \leq \frac{(\gamma^2 - 1)e^\varepsilon}{(\gamma - 1)(e^\varepsilon + \gamma)} \quad P_0\text{-a.e..}$$

Or equivalently,

$$\frac{\gamma + 1}{e^\varepsilon + \gamma} \leq \frac{d\mathbf{Q}_{g_\varepsilon}^*}{dP_0} \leq \frac{(\gamma + 1)e^\varepsilon}{e^\varepsilon + \gamma} \quad P_0\text{-a.e..}$$

This shows that $\mathbf{Q}_{g_\varepsilon}^* \in \tilde{\mathcal{P}}_{b, be^\varepsilon, P_0}$ and completes the proof. □

E Detailed experimental setup and additional experiments

E.1 Mixture of Laplace distributions

Proof of $\tilde{\mathcal{P}}_{\mathcal{L}} \subseteq \tilde{\mathcal{P}}_{e^{-1/b}, e^{1/b}, h_{\mathcal{L}}}$ (Example 3.1)

Let the mixture of Laplace distributions with fixed scale parameter b be defined as

$$\tilde{\mathcal{P}}_{\mathcal{L}} = \left\{ \sum_{i=1}^k \lambda_i \mathcal{L}(m_i, b) : k \in \mathbb{N}, \lambda_i \geq 0, \sum_{i=1}^k \lambda_i = 1, \|m_i\|_1 \leq 1 \right\},$$

where $\mathcal{L}(m, b)$ denotes the n -dimensional Laplace distribution with mean $m \in \mathbb{R}^n$ and scale parameter $b > 0$. We aim to show that

$$\tilde{\mathcal{P}}_{\mathcal{L}} \subseteq \tilde{\mathcal{P}}_{e^{-1/b}, e^{1/b}, h_{\mathcal{L}}},$$

where $h_{\mathcal{L}}$ is the density of the zero-mean n -dimensional Laplace distribution with scale parameter b .

The density of an n -dimensional Laplace distribution with location $m \in \mathbb{R}^n$ and scale b is given by

$$\mathcal{L}(x | m, b) = \frac{1}{(2b)^n} \exp\left(-\frac{\|x - m\|_1}{b}\right), \quad x \in \mathbb{R}^n.$$

Now, fix any $m \in \mathbb{R}^n$ such that $\|m\|_1 \leq 1$. Then for any $x \in \mathbb{R}^n$, we have the following bounds:

$$\begin{aligned} \frac{\|x\|_1 - \|m\|_1}{b} &\leq \frac{\|x - m\|_1}{b} \leq \frac{\|x\|_1 + \|m\|_1}{b} \\ \Rightarrow \frac{\|x\|_1 - 1}{b} &\leq \frac{\|x - m\|_1}{b} \leq \frac{\|x\|_1 + 1}{b}. \end{aligned}$$

Exponentiating both sides, we get:

$$\exp\left(-\frac{\|x\|_1 + 1}{b}\right) \leq \exp\left(-\frac{\|x - m\|_1}{b}\right) \leq \exp\left(-\frac{\|x\|_1 - 1}{b}\right).$$

Multiplying by the constant $\frac{1}{(2b)^n}$, we obtain:

$$\frac{e^{-1/b}}{(2b)^n} \exp\left(-\frac{\|x\|_1}{b}\right) \leq \mathcal{L}(x | m, b) \leq \frac{e^{1/b}}{(2b)^n} \exp\left(-\frac{\|x\|_1}{b}\right),$$

which implies

$$e^{-1/b} \mathcal{L}(x | \mathbf{0}, b) \leq \mathcal{L}(x | m, b) \leq e^{1/b} \mathcal{L}(x | \mathbf{0}, b).$$

For any distribution $P \in \tilde{\mathcal{P}}_{\mathcal{L}}$, we can write

$$p(x) = \sum_{i=1}^k \lambda_i \mathcal{L}(x | m_i, b).$$

Since each m_i satisfies $\|m_i\|_1 \leq 1$, the above inequality applies to each mixture component. Thus, for all $x \in \mathbb{R}^n$,

$$e^{-1/b} \mathcal{L}(x | \mathbf{0}, b) \leq \sum_{i=1}^k \lambda_i \mathcal{L}(x | m_i, b) \leq e^{1/b} \mathcal{L}(x | \mathbf{0}, b).$$

Defining $p_0(x) = \mathcal{L}(x | \mathbf{0}, b)$, we conclude that for every $P \in \tilde{\mathcal{P}}_{\mathcal{L}}$,

$$e^{-1/b} \leq \frac{p(x)}{p_0(x)} \leq e^{1/b}, \quad \forall x \in \mathbb{R}^n,$$

which confirms that $\tilde{\mathcal{P}}_{\mathcal{L}} \subseteq \tilde{\mathcal{P}}_{e^{-1/b}, e^{1/b}, h_{\mathcal{L}}}$.

Experimental details of Figure 1

The original distribution is a mixture of four two-dimensional Laplace distributions, each with scale parameter $b = 2$ and means at $(1, 0)$, $(0, 1)$, $(-1, 0)$, and $(0, -1)$, respectively, with equal weights $\frac{1}{4}$. That is, the input distribution is given by

$$P_{\text{input}} = \sum_{i=1}^4 \frac{1}{4} \mathcal{L}(m_i, 2),$$

where the m_i are defined as above.

From Example 3.1, we know that $P_{\text{input}} \in \tilde{\mathcal{P}}$, where

$$\tilde{\mathcal{P}} = \left\{ \sum_{i=1}^k \lambda_i \mathcal{L}(m_i, b) : k \in \mathbb{N}, \lambda_i \geq 0, \sum_{i=1}^k \lambda_i = 1, \|m_i\|_1 \leq 1 \right\},$$

and furthermore, $\tilde{\mathcal{P}} \subseteq \tilde{\mathcal{P}}_{e^{-1/b}, e^{1/b}, h_{\mathcal{L}}}$, where $h_{\mathcal{L}}$ denotes the density of a two-dimensional Laplace distribution with mean zero and scale parameter $b = 2$.

For $b = 2$, we define the local and global minimax universes as

$$\tilde{\mathcal{P}}_{\text{local}} = \tilde{\mathcal{P}}_{\frac{1}{2}, 2, h_{\mathcal{L}}} \quad \text{and} \quad \tilde{\mathcal{P}}_{\text{global}} = \tilde{\mathcal{P}}_{\frac{1}{6}, 6, h_{\mathcal{L}}}.$$

These universes are chosen such that $\tilde{\mathcal{P}}_{\text{local}} \subseteq \tilde{\mathcal{P}}_{\text{global}}$, $P_{\text{input}} \in \tilde{\mathcal{P}}_{\text{local}} \cap \tilde{\mathcal{P}}_{\text{global}}$, and the ratio $\frac{c_2 - c_1}{1 - c_1} \in \mathbb{N}$, as required by Assumption 3.2. Moreover, we note that $\tilde{\mathcal{P}}_{e^{-1/b}, e^{1/b}, h_{\mathcal{L}}} \subseteq \tilde{\mathcal{P}}_{\frac{1}{2}, 2, h_{\mathcal{L}}}$, i.e., the condition in Assumption 3.2 is achieved by slightly adjusting the bounds c_1 and c_2 using floor and ceiling functions.

We evaluate the performance of minimax-optimal samplers on the input distribution under two LDP settings: comparing local minimax-optimal samplers with global minimax-optimal samplers under both ν -GLDP and ε -LDP. In the pure-LDP setting, we compare the global minimax-optimal sampler of Park et al. [17] with the local minimax-optimal sampler from Theorem 5.1 for $\varepsilon = 1$. In the ν -GLDP setting, we compare our global minimax-optimal sampler from Corollary 3.8 with the local minimax-optimal sampler from Theorem 4.1, using the special case $g = G_{\nu}$ with $\nu = 1.5$.

Under both settings, Figure 1 demonstrates that the local minimax-optimal sampler better preserves the input distribution compared to the global minimax-optimal sampler, given the same level of privacy.

E.2 Gaussian mixtures

Proof of $\tilde{\mathcal{P}}_{\mathcal{N}} \subseteq \tilde{\mathcal{P}}_{0, 1, h_{\mathcal{N}}}$ (Example 3.1)

let $\tilde{\mathcal{P}}_{\mathcal{N}} = \left\{ \sum_{i=1}^k \lambda_i \mathcal{N}(m_i, \sigma^2 I_n) : \lambda_i \geq 0, \sum_{i=1}^k \lambda_i = 1, \|m_i\|_2 \leq 1 \right\}$. We want to show that $\tilde{\mathcal{P}}_{\mathcal{N}} \subseteq \tilde{\mathcal{P}}_{0, 1, h_{\mathcal{N}}}$, where $h_{\mathcal{N}}(x) = \frac{1}{(2\pi\sigma^2)^{\frac{n}{2}}} \exp\left(-\frac{[\max(0, \|x\|_2 - 1)]^2}{2\sigma^2}\right)$.

Fix a component centre $m \in \mathbb{R}^n$ with $\|m\|_2 \leq 1$. Its Gaussian density at x is

$$\mathcal{N}_{m, \sigma^2 I_n}[x] = \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left(-\frac{\|x - m\|_2^2}{2\sigma^2}\right).$$

By the triangle inequality,

$$\|x - m\|_2 \geq \left| \|x\|_2 - \|m\|_2 \right| \geq \max(0, \|x\|_2 - 1),$$

because $\|m\|_2 \leq 1$. Squaring and dividing by $2\sigma^2$ yields

$$\exp\left(-\frac{\|x - m\|_2^2}{2\sigma^2}\right) \leq \exp\left(-\frac{[\max(0, \|x\|_2 - 1)]^2}{2\sigma^2}\right).$$

Multiplying by the common normalizing constant $(2\pi\sigma^2)^{-n/2}$ gives

$$\mathcal{N}_{m, \sigma^2 I_n}[x] \leq h_{\mathcal{N}}(x), \quad \forall x \in \mathbb{R}^n.$$

Now take an arbitrary mixture $P = \sum_{i=1}^k \lambda_i \mathcal{N}(m_i, \sigma^2 I_n) \in \tilde{\mathcal{P}}_{\mathcal{N}}$. Its density is

$$p(x) = \sum_{i=1}^k \lambda_i \mathcal{N}_{m_i, \sigma^2 I_n}[x].$$

Because each component satisfies $\mathcal{N}_{m_i, \sigma^2 I_n}[x] \leq h_{\mathcal{N}}(x)$ and the weights obey $\lambda_i \geq 0$, $\sum_i \lambda_i = 1$, we have

$$0 \leq p(x) \leq h_{\mathcal{N}}(x) \quad \forall x \in \mathbb{R}^n.$$

Therefore $P \in \tilde{\mathcal{P}}_{0,1,h_{\mathcal{N}}}$, as claimed.

It is worth noting that $h_{\mathcal{N}}$ is not a probability density. In order to make $h_{\mathcal{N}}$ a valid probability distribution, we normalize it by its integral and define

$$c_2 = \int h_{\mathcal{N}}(x) dx, \quad \text{and} \quad h(x) = \frac{h_{\mathcal{N}}(x)}{c_2}.$$

Accordingly, we have $\tilde{\mathcal{P}}_{0,1,h_{\mathcal{N}}} = \tilde{\mathcal{P}}_{0,c_2,h}$.

E.3 Finite sample space numerical results under pure LDP

Experimental details of Figure 3

We compare local and global minimax-optimal samplers in the finite setting $\mathcal{X} = [k]$, where the global universe is defined as $\tilde{\mathcal{P}}_{\text{global}} = \mathcal{P}([k]) = \tilde{\mathcal{P}}_{0,k,\mu_k}$, with μ_k denoting the uniform distribution over $[k]$. The local neighborhood is specified as $\tilde{\mathcal{P}}_{\text{local}} = \mathcal{N}_{\gamma}(\mu_k) = \tilde{\mathcal{P}}_{\frac{1}{\gamma},\gamma,\mu_k}$, where $\gamma = \frac{k}{2} - 1$, consistent with the finite-space version of Theorem 5.1.

For the selected values $k \in \{10, 20, 100\}$, it is easy to verify that $\gamma \in \mathbb{N}$, and the uniform distribution μ_k is $(\alpha, \frac{1}{\alpha}, 1)$ -decomposable, where

$$\alpha = \frac{1 - c_1}{c_2 - c_1} = \frac{1 - \frac{1}{\gamma}}{\gamma - \frac{1}{\gamma}} = \frac{1}{\gamma + 1} = \frac{2}{k}.$$

Therefore, the decomposability condition in Theorem 5.1 is satisfied, and we can apply its finite-space version by setting $\mathcal{X} = [k]$.

The local minimax risk is then given by

$$\mathcal{R}(\mathcal{Q}_{[k],\mathcal{P}([k]),\varepsilon}, N_{\gamma}(\mu_k), f) = \frac{1 - r_1}{r_2 - r_1} f(r_2) + \frac{r_2 - 1}{r_2 - r_1} f(r_1),$$

where

$$b := \frac{\gamma - \frac{1}{\gamma}}{(e^{\varepsilon} - 1)(1 - \frac{1}{\gamma}) + \gamma - \frac{1}{\gamma}} = \frac{\gamma + 1}{\gamma + e^{\varepsilon}}, \quad r_1 := \frac{1}{\gamma b} = \frac{e^{\varepsilon} + \gamma}{\gamma(\gamma + 1)}, \quad r_2 := \frac{\gamma}{be^{\varepsilon}} = \frac{\gamma(e^{\varepsilon} + \gamma)}{e^{\varepsilon}(\gamma + 1)}.$$

On the other hand, the global minimax risk is given by [17, Theorem 3.1]:

$$\mathcal{R}(\mathcal{Q}_{[k],\mathcal{P}([k]),\varepsilon}, \mathcal{P}([k]), f) = \frac{e^{\varepsilon}}{e^{\varepsilon} + k - 1} f\left(\frac{e^{\varepsilon} + k - 1}{e^{\varepsilon}}\right) + \frac{k - 1}{e^{\varepsilon} + k - 1} f(0).$$

Therefore, in Figure 3, we compare the theoretical worst-case f -divergence losses achieved by the local minimax-optimal sampler and the global minimax-optimal sampler:

$$\mathcal{R}(\mathcal{Q}_{[k],\mathcal{P}([k]),\varepsilon}, N_{\gamma}(\mu_k), f) \quad \text{and} \quad \mathcal{R}(\mathcal{Q}_{[k],\mathcal{P}([k]),\varepsilon}, \mathcal{P}([k]), f).$$

Additional numerical results

We replicated the procedure used in Section 6.1 to generate Figure 3, but now with sample sizes $k = 10$ and $k = 100$. As shown in Figures 5 and 6 and comparing with Figure 3, the local minimax sampler consistently attains a smaller worst-case loss than the global minimax sampler across different ε values. Also, the performance gap decreases for larger values of k .

E.4 Experimental details of continuous sample space

In the continuous setting with $\mathcal{X} = \mathbb{R}$, we fix the universe $\tilde{\mathcal{P}}_{\text{local}}$ and evaluate the empirical worst-case f -divergence over 100 randomly generated client distributions $P_1, \dots, P_{100} \in \tilde{\mathcal{P}}_{\text{local}}$. Each P_j

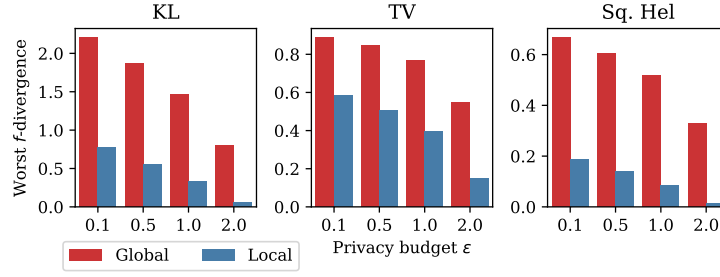


Figure 5: Theoretical worst-case f -divergences of global and local minimax samplers under the pure LDP setting with uniform reference distribution μ_k over finite space ($k = 10$).

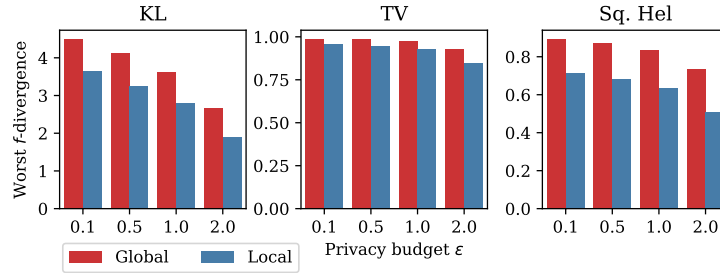


Figure 6: Theoretical worst-case f -divergences of global and local minimax samplers under the pure LDP setting with uniform reference distribution μ_k over finite space ($k = 100$).

represents a client and is constructed as a mixture of a random number of one-dimensional Laplace components with scale parameter $b = 1$. The full generation procedure is described below.

We define the mixture class of Laplace distributions with fixed scale parameter b as

$$\tilde{\mathcal{P}}_{\mathcal{L}} = \left\{ \sum_{i=1}^k \lambda_i \mathcal{L}(m_i, b) : k \in [K], \lambda_i \geq 0, \sum_{i=1}^k \lambda_i = 1, |m_i| \leq 1 \right\},$$

where $\mathcal{L}(m, b)$ denotes the one-dimensional Laplace distribution with mean $m \in \mathbb{R}$ and scale parameter $b > 0$. To prevent an unbounded number of components in each mixture, we impose an upper bound K on the number of Laplace components per client.

Each $P_j \in \tilde{\mathcal{P}}_{\mathcal{L}}$ is generated by randomly selecting k , λ_i , and m_i as follows: First, sample $\tilde{k}$ from a Poisson distribution with mean k_0 , and set $k = \min(\tilde{k} + 1, K)$. Then, sample each $m_1, \dots, m_k$ independently from the uniform distribution on $[-1, 1]$, and sample weights $(\lambda_1, \dots, \lambda_k)$ from the uniform distribution on $\mathcal{P}([k])$. In this experiment, to maintain consistency with Park et al. [17], we use $K = 10$ and $k_0 = 2$.

We define the local and global universes as

$$\tilde{\mathcal{P}}_{\text{local}} = \tilde{\mathcal{P}}_{1/3, 3, h_{\mathcal{L}}} \quad \text{and} \quad \tilde{\mathcal{P}}_{\text{global}} = \tilde{\mathcal{P}}_{1/9, 9, h_{\mathcal{L}}},$$

where $h_{\mathcal{L}}$ denotes the density of the Laplace distribution with mean zero and scale $b = 1$.

We evaluate the empirical worst-case divergence of each sampler using the maximum

$$\max_{j \in [100]} D_f(P_j \| \mathbf{Q}(P_j)).$$

The local minimax sampler is instantiated from Theorem 5.1, while the global minimax sampler corresponds to the optimal sampler from [17, Theorem 3.3].

E.5 Finite sample space numerical results under GLDP

In this section, we adopt the same experimental setup as in Section 6.1 to evaluate the worst-case f -divergence of the local and global minimax samplers under ν -GLDP constraints. To this end, we

follow the same procedure using the same global and local universes. However, the computation of minimax risk differs: for the global minimax risk, we use the optimal value corresponding to the optimal sampler characterized in Corollary 3.8, while for the local minimax risk, we instantiate Theorem 4.1 with $g = G_\nu$ (both the finite sample space version of the results).

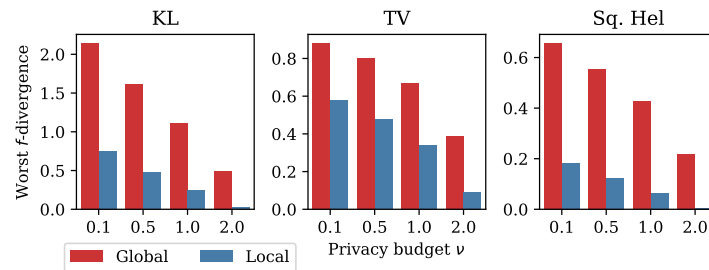


Figure 7: Theoretical worst-case f -divergences of global and local minimax samplers under the ν -GLDP setting with uniform reference distribution μ_k over finite space ($k = 10$).

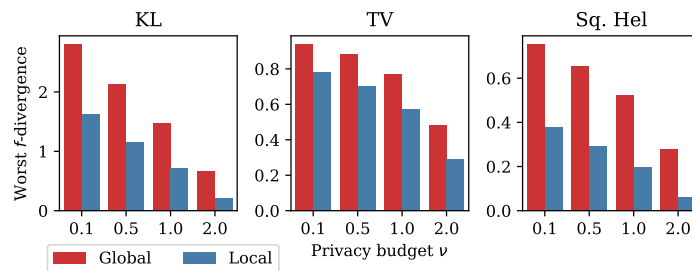


Figure 8: Theoretical worst-case f -divergences of global and local minimax samplers under the ν -GLDP setting with uniform reference distribution μ_k over finite space ($k = 20$).

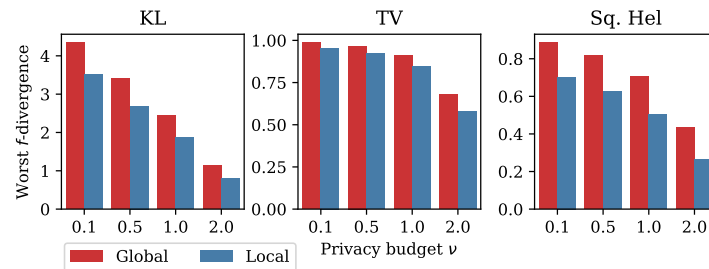


Figure 9: Theoretical worst-case f -divergences of global and local minimax samplers under the ν -GLDP setting with uniform reference distribution μ_k over finite space ($k = 100$).

We compare the worst-case f -divergence under ν -GDP for local and global minimax-optimal samplers across different values of $\nu \in \{0.1, 0.5, 1, 2\}$. The numerical results in Figures 7, 8, and 9 demonstrate that for various sample sizes $k \in \{10, 20, 100\}$, the local minimax-optimal sampler consistently achieves lower minimax risk than the global sampler across all privacy parameter values.

E.6 Continuous sample space numerical results under GLDP

We follow the same experimental procedure described in Appendix E.4, with the only difference being that we now compare the local and global minimax-optimal samplers under ν -GDP instead of ϵ -LDP. The local and global universes, as well as the distribution generation process, remain unchanged. For the global minimax sampler, we use the construction provided in Corollary 3.8, while the local minimax sampler is instantiated from Theorem 4.1 with $g = G_\nu$. We conduct the

comparison across various values of $\nu \in \{0.1, 0.5, 1, 2\}$, evaluating the resulting samplers using three f -divergences: KL divergence, total variation distance, and squared Hellinger distance.

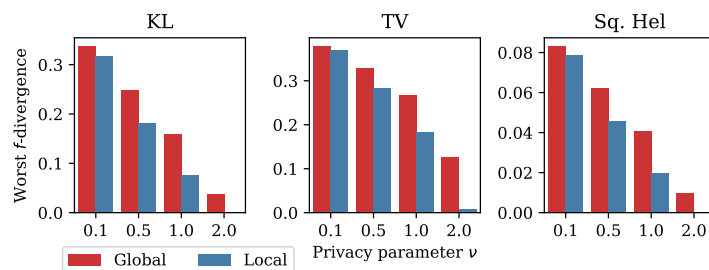


Figure 10: Empirical worst-case f -divergences of global and local minimax samplers under ν -GLDP setting, over 100 experiments on a 1-D Laplace mixture.

As illustrated in Figure 10, in the ν -GDP setting—similar to the ε -LDP case—the local minimax sampler consistently achieves a smaller worst-case f -divergence than the global minimax sampler across all three divergences and privacy parameter values.

E.7 Additional numerical results on continuous sample space

We extend the experiment from Section 6.2 to the two-dimensional setting $\mathcal{X} = \mathbb{R}^2$. Figure 11 compares the worst-case empirical f -divergence between the global sampler of Park et al. [17] and the local sampler described in Theorem 4.1 for $g = g_\varepsilon$.

We follow the same experimental procedure as in Section 6.2, detailed in Appendix E.4, to generate 100 client distributions. Each distribution is constructed as a mixture of 2-D Laplace components with fixed scale parameter $b = 1$. The only difference from the one-dimensional case is the extension to 2-D Laplace mixtures.

The mixture class of 2-D Laplace distributions with scale parameter b is defined as

$$\tilde{\mathcal{P}}_{\mathcal{L}} = \left\{ \sum_{i=1}^k \lambda_i \mathcal{L}(m_i, b) : k \in [K], \lambda_i \geq 0, \sum_{i=1}^k \lambda_i = 1, \|m_i\|_1 \leq 1 \right\},$$

where $\mathcal{L}(m, b)$ denotes a two-dimensional Laplace distribution with mean $m \in \mathbb{R}^2$ and scale $b > 0$. To control the complexity of each mixture, we impose an upper bound K on the number of components per distribution.

Each distribution $P_j \in \tilde{\mathcal{P}}_{\mathcal{L}}$ is generated by randomly sampling k , λ_i , and m_i as follows: First, sample $\tilde{k} \sim \text{Poisson}(k_0)$, and set $k = \min(\tilde{k} + 1, K)$. Then, sample each $m_i \in \mathbb{R}^2$ independently from the uniform distribution on the ℓ_1 ball $\{x \in \mathbb{R}^2 : \|x\|_1 \leq 1\}$, and draw the weights $(\lambda_1, \dots, \lambda_k)$ uniformly from the probability simplex $\mathcal{P}([k])$. In this experiment, following the setup in Park et al. [17], we use $K = 10$ and $k_0 = 2$.

The local and global universes are defined as $\tilde{\mathcal{P}}_{\text{local}} = \tilde{\mathcal{P}}_{1/3, 3, h_{\mathcal{L}}}$ and $\tilde{\mathcal{P}}_{\text{global}} = \tilde{\mathcal{P}}_{1/9, 9, h_{\mathcal{L}}}$, where $h_{\mathcal{L}}$ denotes the density of the two-dimensional Laplace distribution with mean zero and scale parameter $b = 1$.

In the two-dimensional case, similar to the one-dimensional setting, the local sampler outperforms the global sampler for nearly all fixed values of the privacy parameter.

F Instructions for reproducing results

In this appendix, we provide instructions for reproducing the experiments and figures in the paper. For a detailed description of the code, please refer to the provided file `README.md`. For tasks that require substantial runtime, we specify the running times. Tasks that complete in less than 5 seconds are excluded from such reporting.

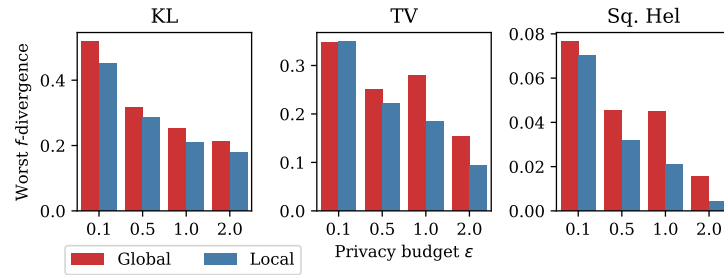


Figure 11: Empirical worst-case f -divergences of global and local minimax samplers under the pure LDP setting, over 100 experiments on a 2-D Laplace mixture.

All experiments were conducted on a system running Ubuntu 22.04.4 LTS, equipped with an Intel(R) Xeon(R) CPU @ 2.20GHz and 16GB of RAM.

F.1 Instructions for reproducing Figure 1

From the repository root, run `python -m experiments.exp_LapMixture_visual` to generate the output distributions, and then `python -m plotting.plot_LapMixture_visual` to produce Figure 1.

The measured running time in our environment for applying the minimax-optimal sampler to the original input distribution is approximately 300 seconds.

F.2 Instructions for reproducing results for finite space

The finite space results are organized into six figures, grouped into two categories. Figures 3, 5, and 6 present the worst-case divergences under pure LDP for different values of $k \in \{10, 20, 100\}$. In contrast, Figures 7, 8, and 9 show the corresponding results under ν -GLDP for the same values of k .

To generate Figures 3, 5, and 6, use the script `plotting/plot_finite_pure.py` with the `-k` argument to specify the desired value of k . For example, the following commands can be used to generate the respective plots:

```
python -m plotting.plot_finite_pure --k 20
python -m plotting.plot_finite_pure --k 10
python -m plotting.plot_finite_pure --k 100
```

To produce Figures 7, 8, and 9, run the script `plotting/plot_finite_GLDP.py` with the corresponding k values:

```
python -m plotting.plot_finite_GLDP --k 10
python -m plotting.plot_finite_GLDP --k 20
python -m plotting.plot_finite_GLDP --k 100
```

F.3 Instructions for reproducing results for continuous space

To reproduce Figure 4, first run the script `experiments/exp_1DLaplaceMix_pure.py` with different values of the privacy parameter ϵ . The following commands correspond to $\epsilon \in \{0.1, 0.5, 1.0, 2.0\}$, respectively:

```
python -m experiments.exp_1DLaplaceMix_pure --eps 0.1 --scale 1 --seed 1
python -m experiments.exp_1DLaplaceMix_pure --eps 0.5 --scale 1 --seed 2
python -m experiments.exp_1DLaplaceMix_pure --eps 1.0 --scale 1 --seed 3
python -m experiments.exp_1DLaplaceMix_pure --eps 2.0 --scale 1 --seed 4
```

These scripts can be executed independently and in any order, or run in parallel. In our environment, running all four in parallel required approximately 600 seconds. After completion, run `python -m plotting.plot_1DLaplaceMix_pure` to generate the final plots.

To reproduce Figure 10, execute the script `experiments/exp_1DLaplaceMix_GLDP.py` with various values of the privacy parameter ν . The following commands correspond to $\nu = 0.1, 0.5, 1.0$, and 2.0 , respectively:

```
python -m experiments.exp_1DLaplaceMix_GLDP --nu 0.1 --scale 1 --seed 1
python -m experiments.exp_1DLaplaceMix_GLDP --nu 0.5 --scale 1 --seed 2
python -m experiments.exp_1DLaplaceMix_GLDP --nu 1.0 --scale 1 --seed 3
python -m experiments.exp_1DLaplaceMix_GLDP --nu 2.0 --scale 1 --seed 4
```

These scripts can be executed in any order or in parallel. In our environment, running all four in parallel required approximately 80 seconds. After completion, run `python -m plotting.plot_1DLaplaceMix_GLDP` to generate the final figure.

To reproduce Figure 11, run the script `experiments/exp_nDLaplaceMix_pure.py` with different values of the privacy parameter ϵ . The following commands correspond to $\epsilon \in \{0.1, 0.5, 1.0, 2.0\}$:

```
python -m experiments.exp_nDLaplaceMix_pure --eps 0.1 --seed 1 --dim 2
python -m experiments.exp_nDLaplaceMix_pure --eps 0.5 --seed 2 --dim 2
python -m experiments.exp_nDLaplaceMix_pure --eps 1.0 --seed 3 --dim 2
python -m experiments.exp_nDLaplaceMix_pure --eps 2.0 --seed 4 --dim 2
```

These can be run independently or in parallel. In our environment, running all four in parallel required approximately 40 seconds. Once the first step is completed, run `python -m plotting.plot_nDLaplaceMix_pure --dim 2` to generate the final figure.

G Broader impact

Our proposed optimal samplers are designed under LDP guarantees and can be applied to privacy protection in generative modeling—a rapidly growing area of interest. A key obstacle to the adoption of privacy-preserving algorithms in practice is the potential degradation in model performance. By characterizing minimax-optimal samplers under two variants of LDP (functional and pure) and across two formulations (global and local), we develop samplers that minimize utility loss under a given privacy budget. As a result, our samplers help overcome this challenge and support broader real-world applicability.

Although an LDP sampler provides meaningful privacy guarantees, it cannot achieve perfect privacy without completely sacrificing utility—there is an inherent trade-off between the two. Moreover, Our samplers are designed for the single-sample setting, where each client releases only one privatized data point. Extending these guarantees to scenarios involving multiple, potentially correlated releases is an important direction for future work (see Section 7). In practice, clients often contribute data through multiple channels, and aggregating these disclosures can lead to greater privacy leakage.