



MAGISTERARBEIT | MASTER'S THESIS

Titel | Title

Local f -Differentially Private Mechanisms and a Case Study on
Efficient Estimation in Bernoulli Models

verfasst von | submitted by

Dr.rer.nat. Annemarie Martha Grass BSc MSc

angestrebter akademischer Grad | in partial fulfilment of the requirements for the degree of
Magistra der Sozial- und Wirtschaftswissenschaften (Mag.rer.soc.oec.)

Wien | Vienna, 2025

Studienkennzahl lt. Studienblatt | Degree
programme code as it appears on the
student record sheet:

UA 066 951

Studienrichtung lt. Studienblatt | Degree
programme as it appears on the student
record sheet:

Magisterstudium Statistik

Betreut von | Supervisor:

Assoz. Prof. Mag. Lukas Steinberger Bakk. BA PhD

Abstract

In the current age, where vast amounts of personal data are constantly collected and analyzed, the need to anonymize potentially sensitive datasets before releasing them to the public is an apparent and pressing problem.

A landmark contribution to this field was the formalization of *differential privacy* (DP) by Dwork and collaborators, which provides a rigorous framework for protecting individual data through randomization. The core idea is to apply a *privacy mechanism* to a dataset, adding controlled noise to ensure that no single individual's data has a significant influence on the output.

This thesis investigates theoretical properties of f -differential privacy (f -DP), a generalization of classical DP grounded in hypothesis testing theory. We provide a sufficient condition under which f -DP mechanisms avoid the undesirable behavior of releasing unsanitized data, a behaviour some classically DP privacy mechanisms can exhibit.

Further, we explore the local variant of f -DP, where users individually apply privacy mechanisms to their own data, e.g. in the absence of a trusted data curator. We study the behaviour of such local mechanisms under aggregation with a special focus on *Gaussian differentially private* mechanisms.

To conclude, we conduct a case study on parameter estimation in Bernoulli models, comparing the different privacy notions in terms of their statistical efficiency.

Kurzzusammenfassung

In der heutigen Zeit werden ständig große Mengen persönlicher Daten gesammelt und analysiert. Das Bedürfnis, potenziell sensible Datensätze vor ihrer Veröffentlichung zu anonymisieren, stellt ein offensichtliches und dringliches Problem dar.

Ein wegweisender Beitrag im dieser Problemstellung zugehörigen Forschungsgebiet war die Formalisierung von *Differential Privacy* (DP) durch Dwork und Koautor*innen. Die zentrale Idee ist, einen sogenannten *Privatisierungsmechanismus* auf den sensiblen Datensatz anzuwenden, um durch kontrolliertes Hinzufügen von Rauschen sicherzustellen, dass einzelne Datenpunkte den Output des Mechanismus nicht wesentlich beeinflussen.

Diese Diplomarbeit untersucht theoretische Eigenschaften von f -Differential Privacy (f -DP), einer Verallgemeinerung des klassischen DP auf Grundlage der Hypothesentesttheorie. Dabei finden wir eine hinreichende Bedingung, unter welcher f -DP-Mechanismen vermeiden, Datensätze unanonymisiert auszugeben - ein unerwünschtes Verhalten, das bei klassischen DP-Mechanismen auftreten kann.

Darüber hinaus widmet sich diese Arbeit der lokalen Variante von f -DP, bei der Nutzer*innen Privatisierungsmechanismen individuell auf ihre Daten anwenden, etwa wenn kein*e vertrauenswürdige*r Datenkurator*in vorhanden ist. Untersucht wird insbesondere das Aggregationsverhalten solcher lokalen Mechanismen, mit besonderem Fokus auf *Gauss'sch*-DP Mechanismen.

Wir schließen mit einer Fallstudie zur Parameterschätzung in Bernoulli-Modellen ab. Die unterschiedlichen Datenschutzkonzepte werden im Hinblick auf ihre statistische Effizienz miteinander verglichen.

Contents

Abstract	i
Kurzzusammenfassung	ii

Thesis

1	Introduction, Setup and Preliminaries	1
2	f-differential privacy and Gaussian differential privacy	2
2.1	The hypothesis test	2
2.2	Gaussian differential privacy	3
2.3	Releasing unsanitized data	4
3	Local f-differential privacy	6
4	Efficient estimation and a case study on Bernoulli data	7
4.1	Identifying optimal mechanisms	8
4.2	Bernoulli data	10
4.3	Comparing optimal mechanisms for different privacy guarantees	11
A	Proof of Proposition 14	16
B	Bernoulli mechanisms	19
B.1	Details on the sanitized model under μ -GDP	19
B.2	Derivation of power-matched privacy parameters	19

1 Introduction, Setup and Preliminaries

In our current age of information and technology, data is now an integral part of nearly every aspect of modern life. The omnipresent smartphones, for instance, continuously collect and store information about their users, ranging from geolocation data and health metrics to web search histories and app usage patterns. These vast and ever-growing datasets are not only central to the commercial strategies of companies but are also increasingly valuable for scientific research and societal planning.

However, such collection and analysis of personal data naturally raises concerns about individual privacy. The need to anonymize datasets is not new and has been recognized long before the current era of “big data”. Unfortunately, early anonymization efforts often relied on ad-hoc methods that proved insufficient in hindsight, as prominently highlighted with the AOL search logs de-anonymization by the New York Times [3], or the Netflix Price de-anonymization [17], among others.

In modern methods, data privacy is achieved by applying randomization to the data before releasing it. This process can be mathematically formalized by applying a *privacy mechanism* to the data, which is represented by a *Markov kernel*. To describe this formally, we consider two measurable spaces: a source space of *sensitive data* $(\mathcal{X}, \mathcal{F})$ and a target space $(\mathcal{Z}, \mathcal{G})$ of *sanitized data*. We think of sanitized data $z \in \mathcal{Z}$ as being generated from sensitive data $x \in \mathcal{X}$ via a privacy mechanism Q , which we represent as a Markov kernel:

Definition 1 (Markov kernel). *A Markov kernel Q is as a mapping $Q : \mathcal{G} \times \mathcal{X} \rightarrow [0, 1]$ satisfying*

- *The map $x \mapsto Q(B|x)$ is $\mathcal{F}$ -measurable for all $B \in \mathcal{G}$.*
- *The map $B \mapsto Q(B|x)$ is a probability measure on $(\mathcal{Z}, \mathcal{G})$ for all $x \in \mathcal{X}$*

A key development in data privacy was the introduction of *differential privacy* (DP) by Dwork [10], which provided a formal framework for quantifying privacy guarantees under randomized data release mechanisms. Throughout this thesis, $n \in \mathbb{N}$ will denote the sample size. We recall the standard definition of ε -differential privacy:

Definition 2 (ε -differential privacy). *A Markov-kernel Q is called ε -differential private (ε -DP) if for all $B \in \mathcal{G}$ and for all $x = (x_1, \dots, x_n), x' = (x'_1, \dots, x'_n) \in \mathcal{X}^n$ with $|\{i : x_i = x'_i\}| = 1$ we have*

$$Q(B|x) \leq e^\varepsilon Q(B|x').$$

Intuitively, differential privacy ensures that the presence or absence of any single individuals data has only a limited effect on the output of the mechanism, controlled by the given parameter ε . Smaller ε ask for closer outputs for neighboring datasets, and thus stronger privacy guarantees, while larger values of ε give weaker privacy guarantees.

This formalization proved widely successful, both in academic research and in industry practice. However, while ε -differential privacy offers a strong and elegant guarantee, its rigidity turned out to make accurate answering of statistical questions difficult unless vast amounts of data are available.

These limitations motivated the introduction of a relaxed variant, known as (ε, δ) -differential privacy, in subsequent work [11, 12]:

Definition 3 $((\varepsilon, \delta)$ -differential privacy). *A Markov-kernel Q is called (ε, δ) -differential private $((\varepsilon, \delta)$ -DP) if for all $B \in \mathcal{G}$ and for all $x, x' \in \mathcal{X}^n$ with $|\{i : x_i = x'_i\}| = 1$ we have*

$$Q(B|x) \leq e^\varepsilon Q(B|x') + \delta.$$

This relaxation permits a small probability δ of the privacy guarantee being violated, allowing mechanisms to achieve improved accuracy while still ensuring a high level of protection in most cases.

The (ε, δ) -definition has since been widely adopted in both research and practice; see [5] for a recent survey. It has also been deployed in real-world systems by major technology companies, including Google [2, 13, 22], Apple [6], Microsoft [7], LinkedIn [18], and the U.S. Census Bureau [1].

However, the fact that $\delta > 0$ corresponds to a nonzero probability of the data being released essentially unsanitized continues to cause discomfort within the privacy community.

These two issues, the overly strict nature of pure differential privacy and the somewhat unsatisfactory interpretation of the δ relaxation, have motivated the investigation of alternative privacy frameworks.

The notion we will focus on in this thesis is *Gaussian differential privacy (GDP)*, introduced in [8]. It builds on the concept of *f-differential privacy (f-DP)*, originally proposed in [21], which is fundamentally based on the following statistical hypothesis testing perspective: Given a privacy mechanism Q and two neighboring datasets D and D' , we consider the binary hypothesis test

$$H_0 : Q \text{ is applied to } D \quad \text{vs.} \quad H_1 : Q \text{ is applied to } D'.$$

Intuitively, this perspective formalizes privacy in terms of how well an adversary could infer which of two neighboring datasets was used to generate a given output. The idea is to model a worst-case attacker who has full knowledge of the mechanism Q and its outputs. The strength of a privacy guarantee can then be understood as a bound on the power of the optimal hypothesis test that such an attacker could perform.

We refer to [20] for a recent survey on this hypothesis testing viewpoint and its impact on the differential privacy community.

Section 2 will formally introduce f -DP, its special case Gaussian differential privacy, and the underlying hypothesis testing framework. In Section 2.3, we provide a sufficient condition under which f -DP mechanisms will not release data unsanitized.

In Section 3 we then turn to the concept of *local* differential privacy, a notion put forward e.g. in [9]. In this setting, no trusted central authority or data curator is assumed, instead each user perturbs their data locally before sending it out, e.g. through a privacy filter on the individual's device or smartphone.

The main focus of this section is to study how local f -DP mechanisms behave under aggregation. In particular we will show that aggregating locally Gaussian differentially private mechanisms again yields a Gaussian differentially private mechanism with the same privacy guarantee.

The thesis concludes with a case study on Bernoulli data in Section 4. There we compare ε -differentially private and Gaussian differentially private mechanisms in terms of their efficiency s on parameter estimation, and investigate whether one framework offers practical advantages over the others.

2 f -differential privacy and Gaussian differential privacy

2.1 The hypothesis test

As mentioned in the introduction, it was first proposed in [21] (and subsequently expanded on in [15]) that differential privacy can be understood in terms of an hypothesis test. The privacy guarantee then corresponds to how well a powerful adversary can distinguish whether the output of a privacy mechanism was generated from one dataset or another. We formalize this idea as follows: Consider two (sensitive) data points $x = (x_1, \dots, x_n), x' = (x'_1, \dots, x'_n) \in \mathcal{X}^n$ with $|\{i : x_i = x'_i\}| = 1$. Given a privacy mechanism $Q : \mathcal{G} \times \mathcal{X} \rightarrow [0, 1]$, we are interested in the question whether our sanitized observation Z was generated from the data point x or x' . This can be formulated as the following hypothesis test:

$$H_0 : Z \sim Q(\cdot | x_1, \dots, x_n) \quad \text{vs.} \quad H_1 : Z \sim Q(\cdot | x'_1, \dots, x'_n) \quad (1)$$

More generally, consider a hypothesis test

$$H_0 : Z \sim \mathbb{P}_0 \quad \text{vs.} \quad H_1 : Z \sim \mathbb{P}_1 \quad (2)$$

Then given a *rejection rule/function* ϕ for this test (2) taking values in $\{0, 1\}$ (or $[0, 1]$) we consider the quantities

$$\begin{aligned}\alpha_\phi &= \mathbb{E}_0[\phi] \\ \beta_\phi &= 1 - \mathbb{E}_1[\phi]\end{aligned}$$

where by $\mathbb{E}_0$ and $\mathbb{E}_1$ we denote the expectations computed under the measures $\mathbb{P}_0$ and $\mathbb{P}_1$ respectively.

Given these quantities we define the *trade-off function* associated to the hypothesis test (2)

$$T(\mathbb{P}_0, \mathbb{P}_1)(\alpha) := \inf_{\phi \text{ is a rejection rule}} \{\beta_\phi : \alpha_\phi \leq \alpha\} \quad (3)$$

It is a direct consequence of the Neyman-Pearson Lemma (cf. [16, Theorem 3.2.1]), that the infimum appearing in the definition of the trade-off function $T(\mathbb{P}_0, \mathbb{P}_1)$ is attained by the likelihood-ratio-test.

Corollary 4. *Consider a general hypothesis test*

$$H_0 : Z \sim \mathbb{P}_0 \quad \text{vs.} \quad H_1 : Z \sim \mathbb{P}_1.$$

Let p_0 and p_1 denote the densities of $\mathbb{P}_0$ and $\mathbb{P}_1$ with respect to a mutual reference measure. For each $\alpha \in [0, 1]$ there exist constants $c \in [0, 1]$ and $k \geq 0$ such that the test defined by

$$\phi(z) := \begin{cases} 1 & \text{if } p_0(z) > k \cdot p_1(z) \\ c & \text{if } p_0(z) = k \cdot p_1(z) \\ 0 & \text{if } p_0(z) < k \cdot p_1(z) \end{cases}$$

enables the following representation of the trade-off function $f := T(\mathbb{P}_0, \mathbb{P}_1)$:

$$f(\alpha) = 1 - \mathbb{E}_1[\phi].$$

The notion of f -differential privacy is defined in terms of the trade-off function associated to the hypothesis test (1).

Definition 5 (f -differential privacy). *Given a trade-off function f , a privacy-mechanism/Markov-kernel Q is said to be f -differentially private, if for all sensitive data points $x, x' \in \mathcal{X}^n$ with $|\{i : x_i = x'_i\}| = 1$ we have*

$$g_Q(\alpha|x, x') := T(Q(\cdot|x), Q(\cdot|x'))(\alpha) \geq f(\alpha). \quad (4)$$

Example 6 ((ϵ, δ) -differential privacy as f -differential privacy). It was observed in [8], that (ϵ, δ) -differentially private mechanism can be seen as f -differentially private mechanism for the choice of

$$f_{(\epsilon, \delta)}(\alpha) := \max\{0, (1 - \delta) - e^\epsilon \alpha, e^{-\epsilon}((1 - \delta) - \alpha)\}. \quad (5)$$

Examples of such trade-off functions can be seen in Figure 1 A.

2.2 Gaussian differential privacy

In the previous section we defined f -differential privacy for general trade-off functions f . To define Gaussian differential privacy we look at the following hypothesis test.

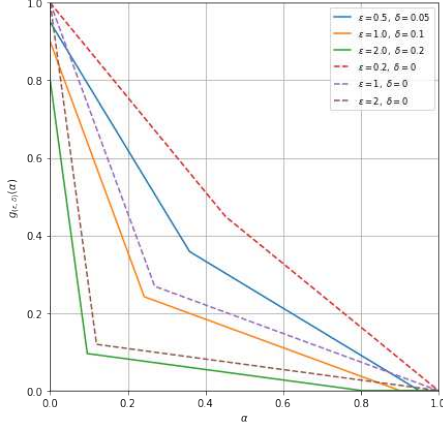
For a given value $\mu \geq 0$ we consider the following test between two normal distributions

$$H_0 : Y \sim \mathcal{N}(0, 1) \quad \text{vs.} \quad H_1 : Y \sim \mathcal{N}(\mu, 1) \quad (6)$$

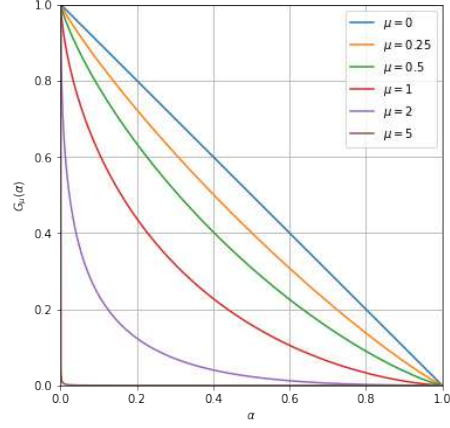
and the associated trade-off function

$$G_\mu(\alpha) := T(\alpha) = T(\mathcal{N}(0, 1), \mathcal{N}(\mu, 1))(\alpha).$$

It is noted in [8] that the following explicit expression for G_μ exists, where Φ denotes the cumulative distribution function of a standard normal.



(A) Examples of the trade-off function $f_{(\varepsilon,\delta)}$ for different values of ε and δ .



(B) Examples of Gaussian trade-off function G_μ for different $\mu \geq 0$.

Figure 1

Proposition 7. *We have that*

$$G_\mu(\alpha) = \Phi(\Phi^{-1}(1 - \alpha) - \mu). \quad (7)$$

Several examples of such Gaussian trade-off functions are depicted in Figure 1 B.

Given this specific trade-off function, we define *Gaussian differential privacy*:

Definition 8 (Gaussian differential privacy). *A privacy-mechanism/Markov-kernel Q is said to be μ -Gaussian differentially private (μ -GDP), if for all sensitive data points $x, x' \in \mathcal{X}^n$ with $|\{i : x_i = x'_i\}| = 1$ we have*

$$T(Q(\cdot|x), Q(\cdot|x'))(\alpha) \geq G_\mu(\alpha). \quad (8)$$

2.3 Releasing unsanitized data

A mechanism satisfying the classic (ε, δ) -DP requirement is known to permit probability δ of releasing the raw data $x \in \mathcal{X}$ unchanged.

We now ask whether or not an f -DP mechanism also shares this often undesirable property.

A sufficient condition for preventing the release of unsanitized data is, that the probability measures induced by the mechanism for any two sensitive datasets are equivalent. Thankfully, for μ -GDP mechanisms this turns out to be the case, and we extend this result to more general f -DP mechanisms:

Proposition 9. *Let $f : [0, 1] \rightarrow [0, 1]$ be a trade-off function satisfying*

$$\begin{aligned} f(0) &= 1, \\ f(\alpha) &\in (0, 1) \text{ for all } \alpha \in (0, 1), \\ f(1) &= 0. \end{aligned} \quad (9)$$

Suppose Q is an f -DP mechanism. Then for any two sensitive datasets $x, x' \in \mathcal{X}^n$, the induced probability measures

$$Q_0 := Q(\cdot|x) \quad \text{and} \quad Q_1 := Q(\cdot|x')$$

are equivalent.

Note that condition (9) is satisfied by the Gaussian trade-off function $G_\mu(\alpha)$ for any $\mu \geq 0$.

Before we are able to give a proof of this proposition, we establish the following key observation.

Proposition 10. *Given two measure $\mathbb{P}_0$ and $\mathbb{P}_1$, consider a general hypothesis test (2). If the associated trade-off function $T(\mathbb{P}_0, \mathbb{P}_1)$ satisfies*

$$T(\mathbb{P}_0, \mathbb{P}_1)(\alpha) > 0 \text{ for all } \alpha \in [0, 1), \quad (10)$$

then (2) admits no (measurable) rejection rule ϕ that has power 1 while maintaining size $\alpha < 1$.

Proof. We proceed by contraposition and assume that the hypothesis test (2), testing $\mathbb{P}_0$ vs. $\mathbb{P}_1$ admits a (measurable) rejection rule $\tilde{\phi}$ with power 1 but size $\tilde{\alpha} < 1$.

Observe the following chain of inequalities

$$1 \geq \sup_{\phi \text{ is a rejection rule}} \{\mathbb{E}_1[\phi] : \mathbb{E}_0[\phi] \leq \tilde{\alpha}\} \geq \mathbb{E}_1[\tilde{\phi}] = 1.$$

For the trade-off function, this implies

$$T(\mathbb{P}_0, \mathbb{P}_1)(\tilde{\alpha}) = 0 \text{ for some } \tilde{\alpha} \in [0, 1),$$

contradicting the assumption in (10). Thus, no such test can exist. $\square$

Proposition 10 enables us to proof that the measures induced by μ -GDP mechanisms are indeed equivalent.

Proof of Proposition 9. Let us first choose $x, x' \in \mathcal{X}^n$ such that $|\{i : x_i = x'_i\}| = 1$. Then

$$T(\mathbb{Q}_0, \mathbb{Q}_1)(\alpha) \geq f(\alpha) \text{ for all } \alpha \in [0, 1],$$

and, in particular, we have that $T(\mathbb{Q}_0, \mathbb{Q}_1)$ satisfies (10)

We first show that $\mathbb{Q}_0 \ll \mathbb{Q}_1$. To this end, assume that there exists a set $A \in \mathcal{G}^{(n)}$ such that $\mathbb{Q}_1(A) = 0$ but $\mathbb{Q}_0(A) = \eta > 0$. Define the following rejection rule

$$\phi_A(z) := \begin{cases} 1 & \text{if } z \notin A, \\ 0 & \text{else.} \end{cases}$$

Then the size of this test is

$$\mathbb{E}_0[\phi_A] = \mathbb{P}_0[Z \notin A] = 1 - \eta < 1$$

while for the power of the test we have

$$\mathbb{E}_1[\phi_A] = \mathbb{P}_1[Z \notin A] = 1.$$

However, Proposition 10 established that such a rejection rule cannot exist for a hypothesis test with a trade-off function that satisfies (10). Therefore, the assumption that such a set A exists leads to a contradiction, and we conclude that $\mathbb{Q}_0$ must be absolutely continuous with respect to $\mathbb{Q}_1$.

As a next step, let us interchange x and x' and observe, that

$$T(\mathbb{Q}_1, \mathbb{Q}_0)(\alpha) = T^{-1}(\mathbb{Q}_0, \mathbb{Q}_1)(\alpha)$$

(cf. [8, Lemma A.2.]). In particular, $T(\mathbb{Q}_1, \mathbb{Q}_0)$ again satisfies (10), and we can analogously conclude $\mathbb{Q}_1 \ll \mathbb{Q}_0$, hence $\mathbb{Q}_0$ and $\mathbb{Q}_1$ are equivalent measures.

Finally, consider $x'' = (x''_1, \dots, x''_n) \in \mathcal{X}^n$ such that $|\{i : x'_i = x''_i\}| = 1$ and $|\{i : x_i = x''_i\}| = 2$. Define $\mathbb{Q}_2 := \mathbb{Q}(\cdot | x'')$. Clearly, we have equivalence of $\mathbb{Q}_1$ and $\mathbb{Q}_2$, hence also of $\mathbb{Q}_0$ and $\mathbb{Q}_2$. The final claim thus follows inductively. $\square$

3 Local f -differential privacy

A key observation in [8] is that Gaussian differential privacy behaves favourably under composition of mechanisms. For each sanitization step, the authors consider (possibly different) mechanisms of the form

$$Q_i : (\mathcal{G}_i, \mathcal{X}^n, \mathcal{Z}_1, \dots, \mathcal{Z}_{i-1}) \rightarrow [0, 1],$$

where the i -th mechanism depends on the full dataset $x \in \mathcal{X}^n$, and the previously sanitized outputs $z_1, \dots, z_{i-1}$. The global aggregated mechanism is then obtained via composition:

$$Q(dz_1, \dots, dz_n | x) := Q_n(dz_n | x; z_{1:n-1}) \cdots Q_i(dz_i | x; z_{1:i-1}) \cdots Q_1(dz_1 | x). \quad (11)$$

The following was shown in [8]: For $i = 1, \dots, n$ let Q_i be a μ_i -GDP mechanisms and define $\boldsymbol{\mu} := (\mu_1, \dots, \mu_n)$. Then the aggregated mechanism Q , as defined in (11), is a ν -GDP mechanism where $\nu := \|\boldsymbol{\mu}\|_2$.

Note that each mechanism Q_i has access to the *full* sensitive dataset $x \in \mathcal{X}^n$ for every sanitization step. We now turn to a *local* paradigm, where each mechanism has access only to a single sensitive data point. For this purpose we adopt the notation from [19]:

Consider measurable spaces $(\mathcal{Z}_1, \mathcal{G}_1), \dots, (\mathcal{Z}_n, \mathcal{G}_n)$, define the aggregated target space $(\mathcal{Z}^{(n)}, \mathcal{G}^{(n)})$ where $\mathcal{Z}^{(n)} := \mathcal{Z}_1 \times \dots \times \mathcal{Z}_n$, and set $\mathcal{G}^{(n)} := \mathcal{G}_1 \otimes \dots \otimes \mathcal{G}_n$.

We consider *sequentially interactive* privacy mechanisms

$$Q_i : (\mathcal{G}_i, \mathcal{X}_i, \mathcal{Z}_i, \dots, \mathcal{Z}_{i-1}) \rightarrow [0, 1] \text{ for } i = 1, \dots, n.$$

Thus, the i -th privatized mechanism has access to the previously privatized data-points $z_1, \dots, z_{i-1}$, as well as the sensitive data point x_i . Note, that this setting naturally includes all *independent* local privacy mechanism, i.e., when Q_i depends only on x_i and does not access previously privatized data. Now, suppose each privacy mechanism is f_i -DP. For this purpose, given two sensitive data $x_i, x'_i \in \mathcal{X}_i$ and previously sanitized data $(z_1, \dots, z_{i-1}) \in \mathcal{Z}^{(i-1)}$, we consider the hypothesis test

$$H_0 : Z_i \sim Q_{i,0} := Q_i(\cdot | x_i; z_{1:i-1}) \quad \text{vs.} \quad H_1 : Z_i \sim Q_{i,1} := Q_i(\cdot | x'_i; z_{1:i-1}). \quad (12)$$

The global aggregated mechanism is given by

$$Q(dz_1, \dots, dz_n | x_1, \dots, x_n) := Q_n(dz_n | x_n; z_{1:n-1}) \cdots Q_i(dz_i | x_i; z_{1:i-1}) \cdots Q_1(dz_1 | x_1). \quad (13)$$

Now, given two sensitive data sets x, x' in $\mathcal{X}^n$ we consider the following global hypothesis test

$$H_0 : Z \sim Q_0 := Q(\cdot | x_1, \dots, x_n) \quad \text{vs.} \quad H_1 : Z \sim Q_1 := Q(\cdot | x'_1, \dots, x'_n). \quad (14)$$

In this setting, we can give a composition theorem that generally strengthens the privacy guarantees given in [8]. To this end, we recall that the *lower convex envelope* $\text{conv}(f)$ of a function $f : [0, 1] \rightarrow [0, 1]$ is given by

$$\text{conv}(f)(\alpha) := \sup \{g(\alpha) | g : [0, 1] \rightarrow \mathbb{R} \text{ convex}, g(\alpha) \leq f(\alpha) \forall \alpha \in [0, 1]\}.$$

Theorem 11. *For $i = 1, \dots, n$ consider sequentially interactive privacy mechanisms*

$$Q_i : (\mathcal{G}_i, \mathcal{X}_i, \mathcal{Z}_i, \dots, \mathcal{Z}_{i-1}) \rightarrow [0, 1] \text{ for } i = 1, \dots, n.$$

Suppose each local mechanism Q_i is f_i -differentially private for some trade-off function $f_i : [0, 1] \rightarrow [0, 1]$.

Then, the aggregated mechanism (13) is f^c -differentially private, where $f^c = \text{conv}(\min\{f_1, \dots, f_n\})$ denotes the lower convex envelope of $\min\{f_1, \dots, f_n\}$.

We can provide even more explicit privacy guarantee for Gaussian-differentially private mechanism:

Corollary 12. *For $i = 1, \dots, n$, consider sequentially interactive privacy mechanisms Q_i , where each Q_i is μ_i -GDP for some $\mu_i \geq 0$. Then, the aggregated mechanism (13) is $\max_{i=1, \dots, n} \mu_i$ -GDP.*

Proof. Note that the Gaussian trade-off function G_μ is monotone in μ , hence

$$\min_{i=1, \dots, n} G_{\mu_i}(\alpha) = G_{\max_{i=1, \dots, n} \mu_i}(\alpha).$$

Moreover, since $G_{\max_{i=1, \dots, n} \mu_i}$ is again a convex function, it follows that

$$\text{conv} \left(\min_{i=1, \dots, n} G_{\mu_i} \right) = G_{\max_{i=1, \dots, n} \mu_i}. \quad \square$$

Remark 13. Corollary 12 highlights that our privacy guarantees are significantly stronger than those given in [8, Corollary 2]. Specifically, for $\boldsymbol{\mu} := (\mu_1, \dots, \mu_n)$, [8, Corollary 2] established that the aggregated mechanism is a $G_{\|\boldsymbol{\mu}\|_2}$ -GDP mechanism. In contrast, Corollary 12 gives that Q is a $G_{\|\boldsymbol{\mu}\|_\infty}$ -GDP mechanism. Since $\|\boldsymbol{\mu}\|_\infty \leq \|\boldsymbol{\mu}\|_2$, our result generally yields a tighter privacy bound. In particular, when all μ_i are equal, i.e. $\mu_1 = \dots = \mu_n$, we obtain $\|\boldsymbol{\mu}\|_\infty = \mu_1$ whereas $\|\boldsymbol{\mu}\|_2 = \sqrt{n}\mu_1$.

The crucial ingredient to the proof of Theorem 11 is the following result. Since its proof is primarily technical, it is deferred to Appendix A.

Proposition 14. *Consider two sensitive datasets $x = (x_1, \dots, x_n), x' = (x'_1, \dots, x'_n) \in \mathcal{X}^n$ such that $|\{i : x_i = x'_i\}| = 1$, and let $i_0 \in \{1, \dots, n\}$ denote the single differing index. Let $g_Q(\cdot|x, x')$ be the trade-off function associated to the hypothesis test (14), and let $g_{i_0}(\cdot|x_{i_0}, x'_{i_0})$ be the trade-off function associated to the respective local hypothesis test (12). Then, for all $\alpha \in [0, 1]$,*

$$g_Q(\alpha|x, x') = g_{i_0}(\alpha|x_{i_0}, x'_{i_0}).$$

Proof of Theorem 11. The goal is to show for any two sensitive data points $x = (x_1, \dots, x_n), x' = (x'_1, \dots, x'_n)$ in $\mathcal{X}^n$ satisfying $|\{i : x_i = x'_i\}| = 1$, the trade-off function of the aggregated mechanism satisfies

$$g_Q(\alpha|x, x') \geq \text{conv}(\min\{f_1, \dots, f_n\})(\alpha) \text{ for all } \alpha \in [0, 1].$$

Thus, let $i_0 \in \{1, \dots, n\}$ denotes the single differing index. By Proposition 14, we have

$$g_Q(\alpha|x, x') = g_{i_0}(\alpha|x_{i_0}, x'_{i_0}) \geq f_{i_0}(\alpha) \text{ for all } \alpha \in [0, 1].$$

In particular, taking the minimum over all indices, this implies

$$g_Q(\alpha|x, x') \geq \min_{i=1, \dots, n} f_i(\alpha) \geq \text{conv}(\min\{f_1, \dots, f_n\})(\alpha) \text{ for all } \alpha \in [0, 1].$$

Since x and x' were arbitrary, we conclude that the global mechanism Q is f^c -differentially private with $f^c = \text{conv}(\min\{f_1, \dots, f_n\})$. $\square$

Remark 15. The proof of Theorem 11 clarifies why the convex envelope appears: the point wise minimum of convex functions is not necessarily convex.

4 Efficient estimation and a case study on Bernoulli data

A long-standing challenge in differential privacy — particularly explored in the framework of ε -differential privacy e.g. in [19] - is the trade-off between strong privacy guarantees and feasibility

of statistical inference. Even in seemingly simple settings, the strictness of ε -DP can severely limit the amount of information available for accurate parameter estimation.

In this section, we identify mechanisms that are optimal for parameter estimation among all mechanisms satisfying the same privacy guarantee.

A key question is whether μ -Gaussian differentially private mechanisms offer advantages over ε -differentially private mechanisms while providing comparable privacy guarantees.

We give a characterization of optimal mechanisms in a general setting, to then deeper investigate data given by a Bernoulli model, where problems become particularly tractable. This enables an in-depth comparison of ε -DP and μ -GDP mechanisms.

4.1 Identifying optimal mechanisms

Let $\Theta \subseteq \mathbb{R}^p$ be a parameter space, and consider a family of statistical models $\mathcal{P} = (P_\theta)_{\theta \in \Theta}$, where each P_θ is a probability measure on a measurable space $(\mathcal{X}, \mathcal{F})$.

Given a model $P_\theta \in \mathcal{P}$ and a privacy mechanism Q , we define the corresponding *sanitized model* by

$$P_\theta^Q := \int_{\mathcal{X}} Q(\cdot | x) P_\theta(dx),$$

and write $Q\mathcal{P} := (P_\theta^Q)_{\theta \in \Theta}$ for the family of sanitized models.

A natural statistical question is to estimate the true parameter θ_0 , but by only using privatized data. This naturally leads to a central question: *Which privacy mechanism Q facilitates the ‘best’ estimation?*

We can formalize this question by considering the *Fisher information* — a well-known measure of how much information a distribution carries about an underlying parameter.

Definition 16. *Given a family of models $\mathcal{P} = (P_\theta)_{\theta \in \Theta}$, assume there exists a common reference measure λ dominating all P_θ , $\theta \in \Theta$. Let $p_\theta := \frac{dP_\theta}{d\lambda}$ denote the corresponding densities, and assume p_θ is differentiable in θ . Then the Fisher information of the model $\mathcal{P}$ at θ is given by*

$$I_\theta(\mathcal{P}) := \mathbb{E}_{X \sim P_\theta} \left[\left(\frac{\nabla_\theta p_\theta(X)}{p_\theta(X)} \right) \left(\frac{\nabla_\theta p_\theta(X)}{p_\theta(X)} \right)^\top \right].$$

(Under suitable regularity conditions, this expression is independent of the choice of dominating measure λ .)

In the same way, one can define the Fisher information of the sanitized model $Q\mathcal{P}$, and compare different mechanisms Q by how much Fisher information they preserve.

A mechanism that enables the most efficient estimation while adhering to a privacy guarantee specified by a trade-off function f can be identified as an optimizer Q^* of the following optimization problem:

$$\sup_{\substack{Q \text{ is an} \\ f\text{-DP mechanism}}} I_\theta(Q\mathcal{P}). \quad (15)$$

Ideally, such an optimizer is independent of the specific choice $\theta \in \Theta$.

Such optimal mechanisms will turn out to be *dominating mechanisms*:

Definition 17. *Let $(\mathcal{Z}^{(1)}, \mathcal{G}^{(1)})$ and $(\mathcal{Z}^{(2)}, \mathcal{G}^{(2)})$ be two sanitized spaces, and let $Q^{(i)} : \mathcal{G}^{(i)} \times \mathcal{X} \rightarrow [0, 1]$ be two privacy mechanisms, for $i \in \{1, 2\}$.*

We say that $Q^{(1)}$ dominates $Q^{(2)}$ if there exists a Markov kernel $\tilde{Q} : \mathcal{G}^{(2)} \times \mathcal{Z}^{(1)} \rightarrow [0, 1]$ such that

$$Q^{(2)}(B | x) = \int_{\mathcal{Z}^{(1)}} \tilde{Q}(B | z) Q^{(1)}(dz | x), \quad \text{for all } x \in \mathcal{X}, B \in \mathcal{G}^{(2)}.$$

In other words, the mechanism $Q^{(1)}$ *dominates* the mechanism $Q^{(2)}$, if $Q^{(2)}$ can be implemented by first applying $Q^{(1)}$, followed by a (possibly randomized) post-processing step $\tilde{Q}$.

The following theorem, due to Blackwell [4] (see also [8, Theorem 2] and [14, Theorem 20]), provides a characterization of dominating mechanisms in terms of their trade-off functions.

Theorem 18. *Let $(\mathcal{Z}^{(1)}, \mathcal{G}^{(1)})$ and $(\mathcal{Z}^{(2)}, \mathcal{G}^{(2)})$ be two sanitized spaces, and let $Q^{(i)} : \mathcal{G}^{(i)} \times \mathcal{X} \rightarrow [0, 1]$ for $i \in \{1, 2\}$ be two privacy mechanisms. Then the following are equivalent:*

(i) *The mechanism $Q^{(1)}$ dominates the mechanism $Q^{(2)}$.*

(ii) *For all $x, x' \in \mathcal{X}^n$, we have*

$$T\left(Q^{(1)}(\cdot | x), Q^{(1)}(\cdot | x')\right)(\alpha) \leq T\left(Q^{(2)}(\cdot | x), Q^{(2)}(\cdot | x')\right)(\alpha) \quad \text{for all } \alpha \in [0, 1].$$

We can now identify mechanisms optimal for your parameter estimation problem.

Theorem 19. *Given a trade-off function f , let*

$$\mathcal{Q}_f := \{Q \text{ is an } f\text{-differentially private mechanism}\}$$

denote the set of all f -DP mechanisms. Then:

(i) *If there exists a mechanism $Q^* \in \mathcal{Q}_f$ that dominates all other mechanisms in $\mathcal{Q}_f$, then*

$$\sup_{Q \in \mathcal{Q}_f} I_\theta(Q\mathcal{P}) = I_\theta(Q^*\mathcal{P}). \quad (16)$$

(ii) *If there exists a mechanism Q_f such that for all sensitive data points $x, x' \in \mathcal{X}^n$ with $|\{i : x_i = x'_i\}| = 1$, we have*

$$T(Q_f(\cdot | x), Q_f(\cdot | x'))(\alpha) = f(\alpha) \quad \text{for all } \alpha \in [0, 1], \quad (17)$$

then

$$\sup_{Q \in \mathcal{Q}_f} I_\theta(Q\mathcal{P}) = I_\theta(Q_f\mathcal{P}).$$

(iii) *If $\mathcal{X} = \{x, x'\}$, then there always exists a mechanism Q_f satisfying (17).*

Proof. The Fisher information is known to satisfy the *data processing inequality*, hence, for any privacy mechanism $Q : \mathcal{G} \times \mathcal{X} \rightarrow [0, 1]$ we have

$$I_\theta(Q\mathcal{P}) \leq I_\theta(\mathcal{P}). \quad (18)$$

In particular, if a mechanism Q is dominated by another mechanism Q^* , then

$$I_\theta(Q\mathcal{P}) \leq I_\theta(Q^*\mathcal{P}),$$

from which (16) clearly follows. Furthermore, statement (ii) clearly follows from (i) and Theorem 18. Finally, the construction of a mechanism Q_f satisfying (17) in the binary setting is given in the proof of [8, Proposition 1]. $\square$

4.2 Bernoulli data

We now turn to the special case $|\mathcal{X}| = 2$, i.e., when the sensitive data can take on only two possible values x and x' . As shown in Theorem 19 (iii), this binary setting is particularly tractable: For any trade-off function f , there always exists an dominating f -differentially mechanism $\mathbb{Q}_f$.

In particular, in this setting, the only hypothesis test that needs to be considered for verifying f -differential privacy is

$$\mathbb{Q}_0 := \mathbb{Q}(\cdot|x) \quad \text{vs.} \quad \mathbb{Q}_1 := \mathbb{Q}(\cdot|x'). \quad (19)$$

Thus, a mechanism $\mathbb{Q}$ is f -differentially private, if and only if for all $\alpha \in [0, 1]$ we have

$$g_{\mathbb{Q}}(\alpha|x, x') := T(\mathbb{Q}_0, \mathbb{Q}_1)(\alpha) \geq f(\alpha). \quad (20)$$

Hence, it suffices to simply prescribe $\mathbb{Q}_0$ and $\mathbb{Q}_1$ in such a way that the resulting trade-off function equals f .

Moreover, for the remainder of this section we assume without loss of generality that $\mathcal{X} = \{0, 1\}$, and on this space of sensitive data we consider a Bernoulli model. That is, we take the parameter space to be $\Theta = [0, 1]$ and define $\mathcal{P}$ as the family of Bernoulli distributions, identified with the probability functions

$$\mathcal{P} = \{p_{\theta} : \theta \in [0, 1]\}, \quad \text{where} \quad p_{\theta}(x) = \theta^x(1 - \theta)^{1-x}, \quad x \in \{0, 1\}.$$

We proceed by providing optimal privacy mechanism in this Bernoulli framework for the privacy notions introduced in Sections 1 and 2. For each privacy guarantee, we present a mechanism whose trade-off function satisfies the respective privacy constraint and thus furthermore maximises the Fisher Information among all such mechanisms.

Gaussian differentially private mechanisms.

Let $\mu \geq 0$, and consider the sanitized space $(\mathcal{Z}, \mathcal{G}) = (\mathbb{R}, \mathcal{B}(\mathbb{R}))$. We define a mechanism $\mathbb{Q}_{\mu} : \mathcal{B}(\mathbb{R}) \times \mathcal{X} \rightarrow [0, 1]$ by

$$\mathbb{Q}_{\mu}(\cdot|0) := \mathcal{N}(0, 1), \quad \mathbb{Q}_{\mu}(\cdot|1) := \mathcal{N}(\mu, 1).$$

The mechanism $\mathbb{Q}_{\mu}$ clearly induces the Gaussian trade-off function

$$g_{\mathbb{Q}_{\mu}}(\alpha|x, x') = G_{\mu}(\alpha), \alpha \in [0, 1],$$

and is therefore a μ -GDP mechanism. It is optimal among all mechanisms satisfying μ -GDP in the sense of maximising the Fisher Information, as guaranteed by Theorem 19 (iii).

The family of sanitized models $\mathbb{Q}_{\mu}\mathcal{P} = \left(P_{\theta}^{\mathbb{Q}}\right)_{\theta \in [0, 1]}$ can be represented by the Gaussian mixtures

$$P_{\theta}^{\mathbb{Q}}(z) = (1 - \theta)\varphi(z) + \theta\varphi(z - \mu),$$

where φ denotes the density of a standard normal distribution. The Fisher Information of this model is given by

$$I_{\theta}(\mathbb{Q}_{\mu}\mathcal{P}) = \int_{\mathbb{R}} \frac{(\varphi(z - \mu) - \varphi(z))^2}{\theta\varphi(z - \mu) + (1 - \theta)\varphi(z)} dz. \quad (21)$$

Further details on this computation are provided in Appendix B.1.

(ϵ, δ)-differentially private mechanisms.

It is the statement of [14, Theorem 18], that if we consider the sanitized data space $\mathcal{Z} := \{0, 1, 2, 3\}$, and define a *quaternary mechanism* $Q_{(\epsilon, \delta)} : \mathcal{B}(\mathcal{Z}) \times \mathcal{X} \rightarrow [0, 1]$ by

$$\begin{aligned} Q_{(\epsilon, \delta)}(0|x) &= \begin{cases} (1-\delta)\frac{1}{1+e^\epsilon} & \text{if } x = 0, \\ (1-\delta)\frac{e^\epsilon}{1+e^\epsilon} & \text{if } x = 1, \end{cases} & Q_{(\epsilon, \delta)}(1|x) &= \begin{cases} (1-\delta)\frac{e^\epsilon}{1+e^\epsilon} & \text{if } x = 0, \\ (1-\delta)\frac{1}{1+e^\epsilon} & \text{if } x = 1, \end{cases} \\ Q_{(\epsilon, \delta)}(2|x) &= \begin{cases} \delta & \text{if } x = 0, \\ 0 & \text{if } x = 1, \end{cases} & Q_{(\epsilon, \delta)}(3|x) &= \begin{cases} 0 & \text{if } x = 0, \\ \delta & \text{if } x = 1. \end{cases} \end{aligned} \quad (22)$$

then the trade-off function is given by

$$g_{Q_{(\epsilon, \delta)}}(\alpha|x, x') = f_{(\epsilon, \delta)}(\alpha) = \max\{0, (1-\delta) - e^\epsilon \alpha, e^{-\epsilon}((1-\delta) - \alpha)\},$$

which coincides with (5).

In particular, when $\delta = 0$, it suffices to consider $\mathcal{Z} = \{0, 1\}$ and the simplified *binary mechanism*

$$Q_\epsilon(0|x) = \begin{cases} \frac{1}{1+e^\epsilon} & \text{if } x = 0, \\ \frac{e^\epsilon}{1+e^\epsilon} & \text{if } x = 1, \end{cases} \quad Q_\epsilon(1|x) = \begin{cases} \frac{e^\epsilon}{1+e^\epsilon} & \text{if } x = 0, \\ \frac{1}{1+e^\epsilon} & \text{if } x = 1. \end{cases} \quad (23)$$

Then the trade-off function becomes

$$g_{Q_\epsilon}(\alpha|x, x') = \begin{cases} 1 - e^\epsilon \alpha & \text{for } \alpha \in \left[0, \frac{1}{e^\epsilon + 1}\right], \\ \frac{1-\alpha}{e^\epsilon} & \text{for } \alpha \in \left(\frac{1}{e^\epsilon + 1}, 1\right]. \end{cases} \quad (24)$$

The associated Fisher Information is given by

$$I_\theta(Q_\epsilon \mathcal{P}) = \frac{(1 - e^\epsilon)^2}{e^\epsilon + \theta(1 - \theta)(1 - e^\epsilon)^2}. \quad (25)$$

4.3 Comparing optimal mechanisms for different privacy guarantees

Which privacy guarantees are more favorable for answering statistical questions? Let us conduct a case-study on Bernoulli models.

Maximally μ -GDP ϵ -DP mechanism

We want to compare the optimal μ -GDP mechanism with the optimal ϵ -DP mechanism in the Bernoulli setting.

To begin, note that for every μ -GDP privacy guarantee, there exists a unique minimal $\epsilon \geq 0$ such that all ϵ -DP mechanisms are also μ -GDP. This value is given by:

$$\epsilon = \log \left(\frac{1 - \Phi\left(-\frac{\mu}{2}\right)}{\Phi\left(-\frac{\mu}{2}\right)} \right). \quad (26)$$

Conversely, for fixed $\epsilon \geq 0$, the binary mechanism Q_ϵ satisfies μ -GDP for any

$$\mu \geq 2\Phi^{-1} \left(\frac{e^\epsilon}{e^\epsilon + 1} \right).$$

In Figure 2 we see how the trade-off functions of the maximally μ -GDP binary privacy mechanism (23) relates to the Gaussian trade-off function. In Figure 3, moreover, we compare the associated Fisher Information. It is, however, clear, that this is not a fair comparison between ϵ -DP and μ -GDP privacy guarantees.

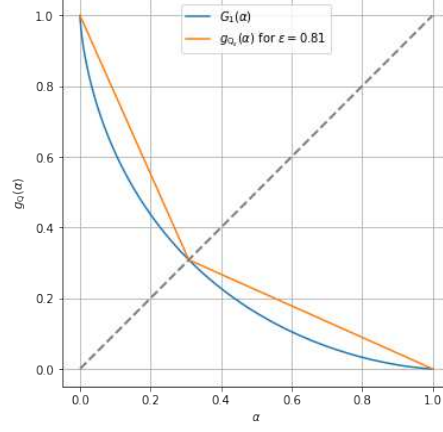


Figure 2: Trade-off functions corresponding to μ -GDP and ε -DP mechanisms are matched in such a way that they touch in α that satisfies $\alpha = G_\mu(\alpha)$.

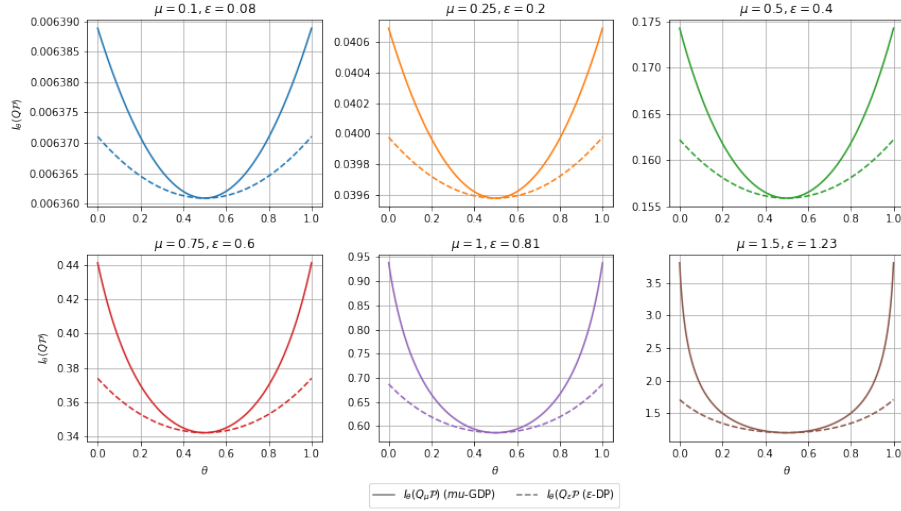


Figure 3: The Fisher Information $I_\theta(QP)$ for the respective privacy mechanism chosen to satisfy Equation (26).

Equally powerful tests

Recall that privacy guarantees can be defined through trade-off functions associated to a hypothesis test. In the binary setting, there is only one trade-off function to consider for the given data, so it becomes natural to build a comparison framework based on these trade-off functions.

That is, we fix a level $\alpha_0 \in [0, 1]$ for the hypothesis test (19), and suppose first that the data is released via a mechanism satisfying a μ -GDP privacy guarantee. We aim to find a corresponding value $\varepsilon_0 \geq 0$ such that the hypothesis test based on an optimal ε_0 -DP mechanism has the same power at level α_0 as the test based on the Gaussian mechanism Q_μ . This requirement is equivalent to solving

$$G_\mu(\alpha_0) = g_{Q_{\varepsilon_0}}(\alpha_0),$$

for ε_0 , where G_μ is the Gaussian trade-off function (7), and $g_{Q_{\varepsilon_0}}$ is the trade-off function (24) associated to the optimal Bernoulli mechanism Q_{ε_0} (23).

One can easily compute, that this $\varepsilon_0 \geq 0$ must be given in the following way

$$\varepsilon = \begin{cases} \log\left(\frac{1-G_\mu(\alpha_0)}{\alpha_0}\right) & \text{if } \alpha_0 \in [0, \Phi(-\frac{\mu}{2})], \\ \log\left(\frac{1-\alpha_0}{G_\mu(\alpha_0)}\right) & \text{if } \alpha_0 \in (\Phi(-\frac{\mu}{2}), 1], \end{cases} \quad (27)$$

the computation can be found in Section B.2.

If we instead fix an ε -DP privacy guarantee with $\varepsilon \geq 0$, then we can identify the corresponding $\mu \geq 0$ such that the Gaussian mechanism yields a test with the same power at a given level $\alpha_0 \in [0, 1]$. This requires choosing

$$\mu = \begin{cases} \Phi^{-1}(e^\varepsilon \alpha) - \Phi^{-1}(\alpha) & \text{if } \alpha \in [0, \frac{1}{e^\varepsilon + 1}], \\ -(\Phi^{-1}(\frac{1-\alpha}{e^\varepsilon}) - \Phi^{-1}(\alpha)) & \text{if } \alpha \in (\frac{1}{e^\varepsilon + 1}, 1], \end{cases} \quad (28)$$

so that the power of the Gaussian test matches that of the binary test. The derivation of this expression is given in Appendix B.2.

In Figure 4 and 5 we give some examples of level-matched trade-off functions. The level for the hypothesis test is fixed at $\alpha_0 = 0.1$ in Figure 4 and at $\alpha_0 = 0.2$ in Figure 5. Each sub-plot shows the binary trade-off function for a different privacy guarantee $\varepsilon > 0$, and the Gaussian trade-off function that matches the power for the given level.

Moreover, we can consider the Fisher Informations associated to the respective optimal mechanisms, computed as in (25) and (21) respectively. Figure 6 gives the Fisher information where the respective privacy guarantees are matched for the level $\alpha_0 = 0.1$, while in Figure 7 they are matched for the level $\alpha_0 = 0.2$.

We observe that there is no universally optimal choice of privacy mechanism for recovering the parameter θ . Rather, the performance seems to depend on the interaction between the privacy parameters α_0 , ε , or μ , and the underlying parameter θ . Our case-study suggest that the μ -GDP mechanism tends to perform better for larger values of ε or μ —that is, in lower-privacy regimes—while the ε -DP mechanism appears more effective under stronger privacy guarantees. A promising direction for future research is to further investigate this trade-off to develop a deeper understanding of the conditions under which each mechanism is preferable.

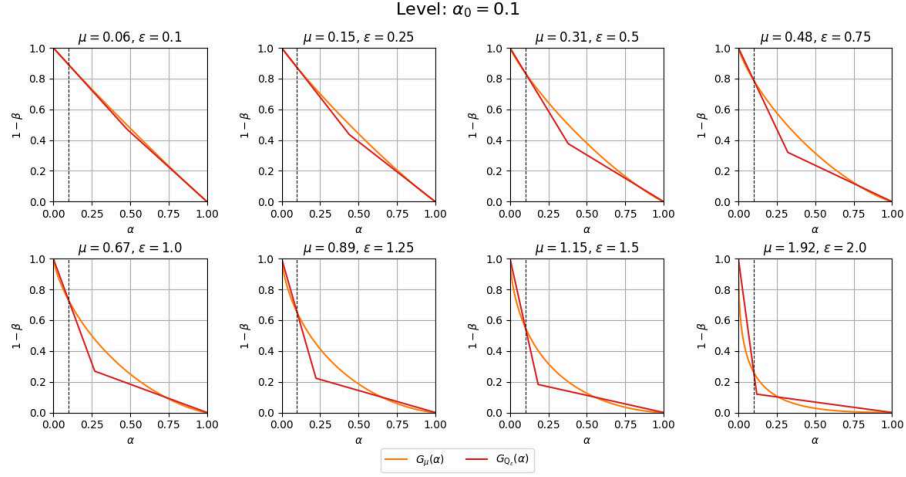


Figure 4: Comparison of trade-off functions for different fixed values of ε . Each subplot corresponds to a different value of $\varepsilon \in [0.1, 2.0]$, and μ is computed such that the trade-off functions intersect at the given level $\alpha_0 = 0.1$.

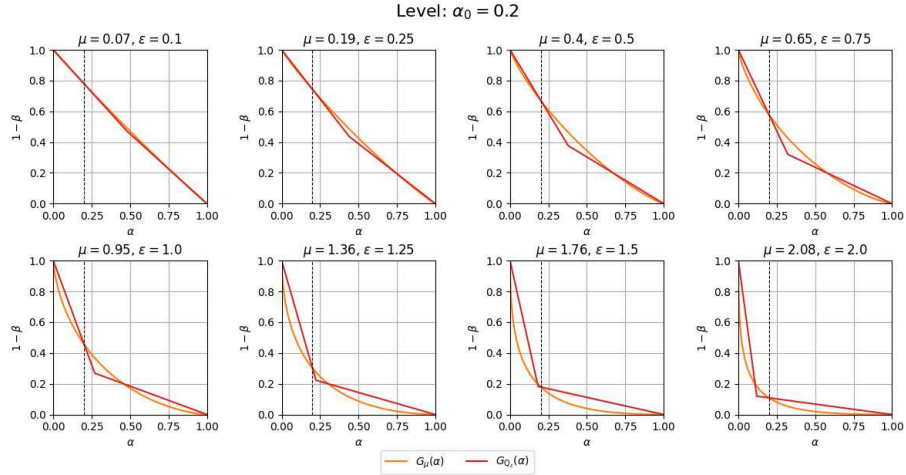


Figure 5: Comparison of trade-off functions for different fixed values of ε . Each subplot corresponds to a different value of $\varepsilon \in [0.1, 2.0]$, and μ is computed such that the trade-off functions intersect at the given level $\alpha_0 = 0.2$.

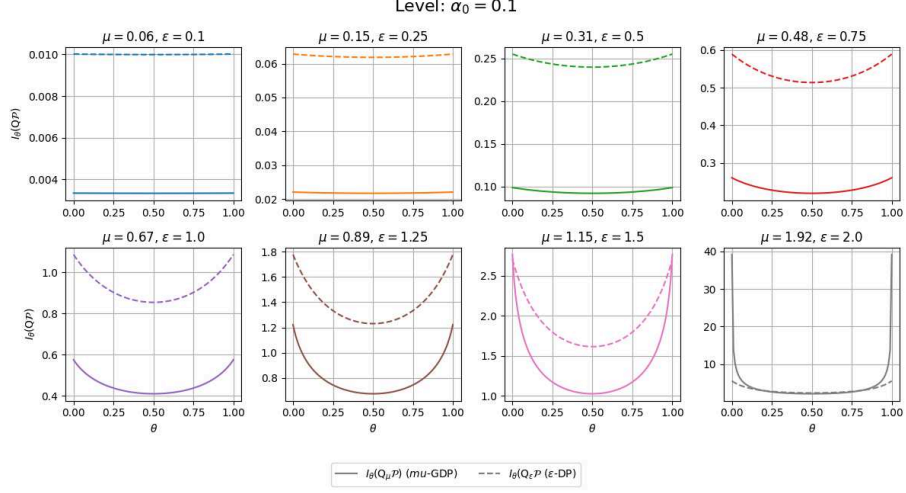


Figure 6: Comparison of Fisher Information under equally powerful privacy mechanisms for the level $\alpha_0 = 0.1$.

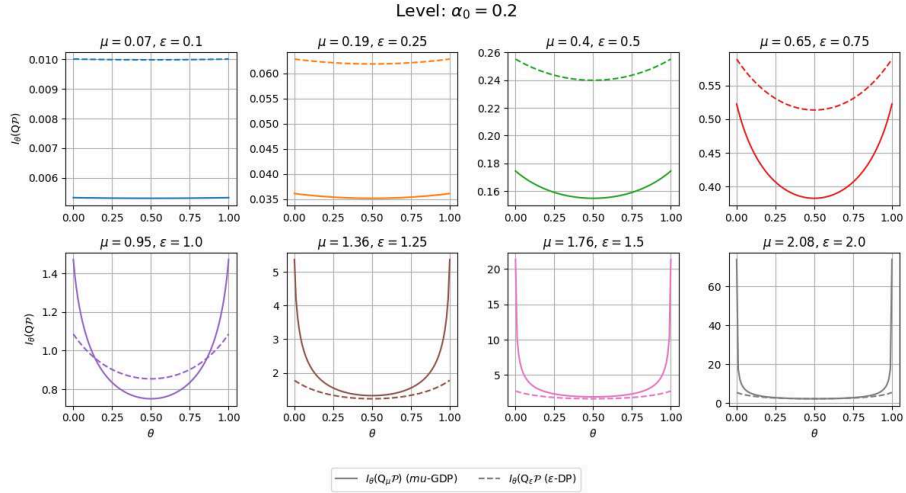


Figure 7: Comparison of Fisher Information under equally powerful privacy mechanisms for the level $\alpha_0 = 0.2$.

A Proof of Proposition 14

The key to proving Proposition 14 lies in representing the measures induced by the aggregated mechanism using suitable density functions. This approach then enables an application of the Neyman Pearson Lemma.

To this end, let x, x' in $\mathcal{X}^n$ denote two sensitive data sets, and let $\mathbb{Q}_0$ and $\mathbb{Q}_1$ respectively denote the measures induced by the (global, aggregated) mechanism $\mathbb{Q}$. The first step is to identify an appropriate common reference measure λ for $\mathbb{Q}_0$ and $\mathbb{Q}_1$, and density functions $\mathbf{q}_0$ and $\mathbf{q}_1$ such that we have the following representations

$$\mathbb{Q}_0(A) = \int_A \mathbf{q}_0(z) \lambda(dz) \quad \text{and} \quad \mathbb{Q}_1(A) = \int_A \mathbf{q}_1(z) \lambda(dz).$$

Furthermore, we want these densities functions to factorize into densities of the respective local mechanisms. For this purpose, for the i -th local mechanism define the reference measure

$$\lambda_i(\cdot | x_i, x'_i; z_{1:i-1}) := \mathbb{Q}_i(\cdot | x_i; z_{1:i-1}) + \mathbb{Q}_i(\cdot | x'_i; z_{1:i-1}).$$

Then clearly

$$\mathbb{Q}_i(\cdot | x_i; z_{1:i-1}), \mathbb{Q}_i(\cdot | x'_i; z_{1:i-1}) \ll \lambda_i(\cdot | x_i, x'_i; z_{1:i-1}),$$

allowing us to apply the Radon-Nikodym Theorem to obtain the (conditional) densities

$$\begin{aligned} q_i(z | x_i, x'_i; z_{1:i-1}) &:= \frac{d\mathbb{Q}_i(\cdot | x_i; z_{1:i-1})}{d\lambda_i(\cdot | x_i, x'_i; z_{1:i-1})}(z), \text{ and} \\ q'_i(z | x_i, x'_i; z_{1:i-1}) &:= \frac{d\mathbb{Q}_i(\cdot | x'_i; z_{1:i-1})}{d\lambda_i(\cdot | x_i, x'_i; z_{1:i-1})}(z). \end{aligned} \tag{29}$$

For every $A_i \in \mathcal{G}_i$, and every $x'_i \in \mathcal{X}_i$ we thus have

$$\mathbb{Q}_{i,0}(A_i) = \mathbb{Q}_i(A_i | x_i; z_{1:i-1}) = \int_{A_i} q_i(z_i | x_i, x'_i; z_{1:i-1}) \lambda_i(dz_i | x_i, x'_i; z_{1:i-1}),$$

while for every $x_i \in \mathcal{X}_i$ we have

$$\mathbb{Q}_{i,1}(A_i) = \mathbb{Q}_i(A_i | x'_i; z_{1:i-1}) = \int_{A_i} q'_i(z_i | x_i, x'_i; z_{1:i-1}) \lambda_i(dz_i | x_i, x'_i; z_{1:i-1}).$$

Given these local densities, it remains to establish a density representation for the global measures $\mathbb{Q}_0$ and $\mathbb{Q}_1$. This follows naturally from the sequential factorization of the mechanism.

Proposition 20. *Given two sensitive data sets $x = (x_1, \dots, x_n), x' = (x'_1, \dots, x'_n) \in \mathcal{X}^n$ let q_i (resp. q'_i) denote the (conditional) density of the measure $\mathbb{Q}_{i,0}$ (resp. $\mathbb{Q}_{i,1}$) as given in (29). Then there exists a common reference measure λ such that*

$$\begin{aligned} \mathbf{q}_0(z_{1:n} | x_{1:n}, x'_{1:n}) &:= q_n(z_n | x_n, x'_n; z_{1:n-1}) \cdots q_1(z_1 | x_1, x'_1; z_1), \text{ and} \\ \mathbf{q}_1(z_{1:n} | x_{1:n}, x'_{1:n}) &:= q'_n(z_n | x_n, x'_n; z_{1:n-1}) \cdots q'_1(z_1 | x_1, x'_1; z_1) \end{aligned}$$

is a density for $\mathbb{Q}_0$ and $\mathbb{Q}_1$ respectively.

Proof. We consider sets of the form $A = \bigotimes_{i=1}^n A_i$, for $A_i \in \mathcal{G}_i$. Then by the definition of $\mathbb{Q}_0$, we have

$$\mathbb{Q}_0(A) = \int_A \mathbb{Q}(dz_1, \dots, dz_n | x_{1:n}) = \int_{A_1} \cdots \int_{A_n} \mathbb{Q}_n(dz_n | x_n; z_{1:n-1}) \cdots \mathbb{Q}_1(dz_1 | x_1).$$

For the innermost integral, we observe

$$\int_{A_n} Q_n(dz_n | x_n; z_{1:n-1}) = \int_{A_n} q_n(z_n | x_n, x'_n; z_{1:n-1}) \lambda_n(dz_n | x_n, x'_n; z_{1:n-1}).$$

Defining this integral as a function of z_{n-1} , we introduce

$$\begin{aligned} h_{n-1}(z_{n-1}) &= h_{n-1}(z_{n-1} | A_n; x_n, x'_n; z_{1:n-2}) \\ &:= \int_{A_n} q_n(z_n | x_n, x'_n; z_{1:n-1}) \lambda_n(dz_n | x_n, x'_n; z_{1:n-1}). \end{aligned}$$

Since h_{n-1} is non-negative and bounded by 1, it is integrable with respect to the probability measure $Q_{n-1}(\cdot | x_{n-1}; z_{1:n-2})$, which leads to

$$\begin{aligned} \int_{A_{n-1}} h_{n-1}(z_{n-1} | A_n; x_n, x'_n; z_{1:n-2}) Q_{n-1}(dz_{n-1} | x_{n-1}; z_{1:n-2}) \\ = \int_{A_{n-1}} h_{n-1}(z_{n-1} | A_n; x_n, x'_n; z_{1:n-2}) \\ q_{n-1}(z_{n-1} | x_{n-1}, x'_{n-1}; z_{1:n-2}) \lambda_{n-1}(dz_{n-1} | x_{n-1}, x'_{n-1}; z_{1:n-2}) \end{aligned}$$

By iterating this argument for all indices, we obtain

$$\begin{aligned} Q_0(A) &= \int_{A_1} \cdots \int_{A_{n-1}} \int_{A_n} q_n(z_n | x_n, x'_n; z_{1:n-1}) \lambda_n(dz_n | x_n, x'_n; z_{1:N-1}) \\ &\quad q_{n-1}(z_{n-1} | x_{n-1}, x'_{n-1}; z_{1:n-2}) \lambda_{n-1}(dz_{n-1} | x_{n-1}, x'_{n-1}; z_{1:n-2}) \\ &\quad \vdots \\ &\quad q_i(z_i | x_i, x'_i; z_{1:i-1}) \lambda_i(dz_i | x_i, x'_i; z_{1:i-1}) \\ &\quad \vdots \\ &\quad q_1(z_1 | x_1, x'_1; z_1) \lambda_1(dz_1 | x_1, x'_1) \\ &= \int_{A_1} \cdots \int_{A_{n-1}} \int_{A_n} \underbrace{q_n(z_n | x_n, x'_n; z_{1:n-1}) \cdots q_1(z_1 | x_1, x'_1; z_1)}_{=: \mathbf{q}_0(z_{1:n} | x_{1:n}, x'_{1:n})} \lambda(dz_1, \dots, dz_n | x_{1:n}, x'_{1:n}) \end{aligned}$$

for

$$\lambda(dz_1, \dots, dz_n | x_{1:n}, x'_{1:n}) := \lambda_n(dz_n | x_n, x'_n; z_{1:N-1}) \cdots \lambda_i(dz_i | x_i, x'_i; z_{1:i-1}) \cdots \lambda_1(dz_1 | x_1, x'_1).$$

Since sets of the form $A = \bigotimes_{i=1}^n A_i$ with $A_i \in \mathcal{G}_i$ generate the product σ -algebra, it follows that the function

$$\mathbf{q}_0(z_{1:n} | x_{1:n}, x'_{1:n}) = q_n(z_n | x_n, x'_n; z_{1:n-1}) \cdots q_1(z_1 | x_1, x'_1; z_1)$$

indeed serves as a density of the measure Q_0 with respect to the reference measure λ .

By an entirely analogous argument, we obtain that

$$\mathbf{q}_1(z_{1:n} | x_{1:n}, x'_{1:n}) = q'_n(z_n | x_n, x'_n; z_{1:n-1}) \cdots q'_1(z_1 | x_1, x'_1; z_1)$$

is a density for Q_1 with respect to the same reference measure λ . □

We can now proceed with the proof of Proposition 14.

Proof of Proposition 14. We begin with the following observation:

For $(z_1, \dots, z_{i-1}) \in \mathcal{Z}^{(i)}$ and $x_i = x'_i \in \mathcal{X}_i$, we clearly have

$$\mathbb{Q}_{i,0} = \mathbb{Q}_i(\cdot | x_i; z_{1:i-1}) = \mathbb{Q}_i(\cdot | x'_i; z_{1:i-1}) = \mathbb{Q}_{i,1}.$$

In particular, this implies that the Radon-Nikodym densities with respect to the common reference measure $\lambda_i = \mathbb{Q}_{i,0} + \mathbb{Q}_{i,1} = 2\mathbb{Q}_i(\cdot | x_i; z_{1:i-1})$, satisfy

$$q_i(z_i) = q_i(z_i | x_i, x'_i; z_{1:i-1}) = q'_i(z_i | x_i, x'_i; z_{1:i-1}) = q'_i(z_i).$$

Now, given two sensitive samples $x = (x_1, \dots, x_n), x' = (x'_1, \dots, x'_n) \in \mathcal{X}^n$ that differ in exactly one coordinate, i.e., $|\{i : x_i = x'_i\}| = 1$, identify the single differing coordinate $i_0 \in \{0, \dots, n\}$. Then we obtain

$$\frac{\mathbf{q}_0(z_{1:n})}{\mathbf{q}_1(z_{1:n})} = \frac{q_n(z_n | x_n, x'_n; z_{1:n-1}) \cdots q_1(z_1 | x_1, x'_1; z_1)}{q'_n(z_n | x_n, x'_n; z_{1:n-1}) \cdots q'_1(z_1 | x_1, x'_1; z_1)} = \frac{q_i(z_i | x_i, x'_i; z_{1:i-1})}{q'_i(z_i | x_i, x'_i; z_{1:i-1})}. \quad (30)$$

Now, given a significance level $\alpha \in [0, 1]$, let ϕ denote the likelihood ratio test associated to the hypothesis test (14) from Corollary 4, with associated constants c, k . Similarly, let ϕ_{i_0} be the likelihood ratio test associated with the hypothesis test (12) with constants c_{i_0} and k_{i_0} .

To obtain equality of the trade-off functions, it suffices, due to Corollary 4, to show that

$$\mathbb{E}_{\mathbb{Q}_1}[\phi] = \mathbb{E}_{\mathbb{Q}_{i_0,1}}[\phi_{i_0}]. \quad (31)$$

This equality is equivalent to the two equalities:

$$\mathbb{Q}_1[\mathbf{q}_1(Z) > k \cdot \mathbf{q}_0(Z)] = \mathbb{Q}_{i_0,1}[q'_{i_0}(Z_{i_0}) > k_{i_0} \cdot q_{i_0}(Z_{i_0})], \quad (32)$$

and

$$c \cdot \mathbb{Q}_1[\mathbf{q}_1(Z) = k \cdot \mathbf{q}_0(Z)] = c_{i_0} \cdot \mathbb{Q}_{i_0,1}[q'_{i_0}(Z_{i_0}) = k_{i_0} \cdot q_{i_0}(Z_{i_0})]. \quad (33)$$

Note, that (30) implies the following relation

$$\left\{z \in \mathcal{Z}^{(n)} : \mathbf{q}_1(z) > k \cdot \mathbf{q}_0(z)\right\} = \left(\bigotimes_{\substack{i=1 \\ i \neq i_0}}^n \mathcal{Z}_i\right) \times \left\{z_{i_0} \in \mathcal{Z}_{i_0} : q'_{i_0}(z_{i_0}) > k \cdot q_{i_0}(z_{i_0})\right\}. \quad (34)$$

This immediately implies, that

$$\mathbb{Q}_1[\mathbf{q}_1(Z) > k \cdot \mathbf{q}_0(Z)] = \mathbb{Q}_{i_0,1}[q'_{i_0}(Z_{i_0}) > k \cdot q_{i_0}(Z_{i_0})].$$

A similar argument holds when replacing inequalities with equalities, leading directly to the conclusion that proving (32) and (33) reduces to showing $c = c_{i_0}$ and $k = k_{i_0}$.

We take from the proof of the Neyman-Pearson Lemma in [16, Theorem 3.2.1], that c and k will be chosen with the help of the functions

$$\begin{aligned} \alpha(k) &= \mathbb{Q}_0[\mathbf{q}_1(Z) > k \cdot \mathbf{q}_0(Z)], \text{ and} \\ \alpha_{i_0}(k) &= \mathbb{Q}_{i_0,0}[q'_{i_0}(Z_{i_0}) > k \cdot q_{i_0}(Z_{i_0})] \end{aligned}$$

respectively. Then k is chosen such that $\alpha(k) \leq \alpha \leq \lim_{\epsilon \searrow 0} \alpha(k - \epsilon)$, and given $\tilde{\alpha} := \lim_{\epsilon \searrow 0} \alpha(k + \epsilon)$, c is defined in the following way

$$c := \frac{\alpha - \alpha(k)}{\tilde{\alpha} - \alpha(k)}.$$

A similar definition holds for c_{i_0} and k_{i_0} . However, as a consequence of (34), we have

$$\alpha(k) = \mathbb{Q}_0[\mathbf{q}_1(Z) > k \cdot \mathbf{q}_0(Z)] = \mathbb{Q}_{i_0,0}[q'_{i_0}(Z_{i_0}) > k \cdot q_{i_0}(Z_{i_0})] = \alpha_{i_0}(k),$$

which concludes the proof of (31). $\square$

B Bernoulli mechanisms

B.1 Details on the sanitized model under μ -GDP

We derive the expression for the sanitized model induced by the mechanism Q_μ , and compute its Fisher Information.

Recall the sensitive space $\mathcal{X} = \{0, 1\}$, and that for $\mu \geq 0$ the Gaussian privacy mechanism $Q_\mu : \mathcal{B}(\mathbb{R}) \times \mathcal{X} \rightarrow [0, 1]$ was given by

$$Q_\mu(\cdot \mid 0) = \mathcal{N}(0, 1), \quad \text{and} \quad Q_\mu(\cdot \mid 1) = \mathcal{N}(\mu, 1).$$

We can represent a privatized model by its density p_θ , which can be derived in the following way. For $z \in \mathcal{Z} = \mathbb{R}$ we have

$$\begin{aligned} p_\theta^Q(z) &= \mathbb{P}_\theta[Z = dz] \\ &= \mathbb{P}_\theta[x = 0]\mathbb{P}_\theta[Z = dz \mid x = 0] + \mathbb{P}_\theta[x = 1]\mathbb{P}_\theta[Z = dz \mid x = 1] \\ &= (1 - \theta)\varphi(z) + \theta\varphi(z - \mu), \end{aligned}$$

where φ denotes the density of a standard normal distribution.

Differentiating with respect to θ gives:

$$\frac{\partial}{\partial \theta} p_\theta^Q(z) = \varphi(z - \mu) - \varphi(z).$$

The Fisher Information of the sanitized model $Q_\mu \mathcal{P}$ is then given by

$$\begin{aligned} I_\theta(Q_\mu \mathcal{P}) &= \mathbb{E}_\theta \left[\left(\frac{\partial}{\partial \theta} \log p_\theta^Q(Z) \right)^2 \right] = \int_{\mathbb{R}} \left(\frac{\varphi(z - \mu) - \varphi(z)}{p_\theta^Q(z)} \right)^2 p_\theta^Q(z) dz \\ &= \int_{\mathbb{R}} \frac{(\varphi(z - \mu) - \varphi(z))^2}{(1 - \theta)\varphi(z) + \theta\varphi(z - \mu)} dz. \end{aligned}$$

To the best of our knowledge, there is no closed-form expression for this integral.

B.2 Derivation of power-matched privacy parameters

We provide details on computing an ε -DP guarantee that yields a hypothesis test with the same power as that under a fixed μ -GDP guarantee, for a given level $\alpha_0 \in [0, 1]$.

Recall that for $\mu \geq 0$, the trade-off function associated with the Gaussian mechanism is given by

$$G_\mu(\alpha) = \Phi(\Phi^{-1}(1 - \alpha) - \mu) = \Phi(-\Phi^{-1}(\alpha) - \mu),$$

where Φ denotes the standard normal cumulative distribution function.

It is easy to check that this function satisfies

$$G_\mu(\tilde{\alpha}) = \tilde{\alpha} \quad \text{for } \tilde{\alpha} := \Phi\left(-\frac{\mu}{2}\right).$$

The binary trade-off function associated with an ε -DP guarantee is piecewise affine:

$$g_{Q_\varepsilon}(\alpha) = \begin{cases} 1 - e^\varepsilon \alpha & \text{for } \alpha \in \left[0, \frac{1}{1+e^\varepsilon}\right], \\ \frac{1 - \alpha}{e^\varepsilon} & \text{for } \alpha \in \left(\frac{1}{1+e^\varepsilon}, 1\right]. \end{cases}$$

To determine the value of ε such that the binary test has the same power as the Gaussian test at level α_0 , we equate the trade-off functions and solve for ε .

-
- If $\alpha_0 \in [0, \Phi(-\frac{\mu}{2})]$, then the relevant part of the binary trade-off function is $1 - e^\varepsilon \alpha_0$. We solve

$$G_\mu(\alpha_0) = 1 - e^\varepsilon \alpha_0 \quad \Leftrightarrow \quad \varepsilon = \log \left(\frac{1 - G_\mu(\alpha_0)}{\alpha_0} \right).$$

- If $\alpha_0 \in (\Phi(-\frac{\mu}{2}), 1]$, then the relevant part is $\frac{1 - \alpha_0}{e^\varepsilon}$. We solve

$$G_\mu(\alpha_0) = \frac{1 - \alpha_0}{e^\varepsilon} \quad \Leftrightarrow \quad \varepsilon = \log \left(\frac{1 - \alpha_0}{G_\mu(\alpha_0)} \right).$$

These expressions yield the formula for the power-matched ε -DP guarantee found in (27).

Conversely, to match the power when ε is fixed, we solve

$$g_{Q_\varepsilon}(\alpha_0) = G_\mu(\alpha_0)$$

depending on α_0 in relation to $\frac{1}{1+e^\varepsilon}$. Thus

- For $\alpha_0 \in [0, \frac{1}{1+e^\varepsilon}]$ we solve

$$\Phi(-\Phi^{-1}(\alpha_0) - \mu) = 1 - e^\varepsilon \alpha_0 \quad \Leftrightarrow \quad \mu = \Phi^{-1}(e^\varepsilon \alpha_0) - \Phi^{-1}(\alpha_0).$$

- For $\alpha_0 \in (\frac{1}{1+e^\varepsilon}, 1]$ we solve

$$\Phi(-\Phi^{-1}(\alpha_0) - \mu) = \frac{1 - \alpha_0}{e^\varepsilon} \quad \Leftrightarrow \quad \mu = - \left(\Phi^{-1} \left(\frac{1 - \alpha_0}{e^\varepsilon} \right) - \Phi^{-1}(\alpha_0) \right).$$

This yields the expression for μ in terms of ε and α_0 stated in Equation (28).

References

- [1] J. M. Abowd and M. B. Hawes. Confidentiality protection in the 2020 US Census of Population and Housing. *Annu. Rev. Stat. Appl.* 10 (2023), pp. 119–144.
- [2] A. Aktay, S. Bavadekar, G. Cossoul, J. Davis, D. Desfontaines, A. Fabrikant, E. Gabrilovich, K. Gadepalli, B. Gipson, M. Guevara, C. Kamath, M. Kansal, A. Lange, C. Mandayam, A. Oplinger, C. Pluntke, T. Roessler, A. Schlosberg, T. Shekel, S. Vispute, M. Vu, G. Wellenius, B. Williams, and R. J. Wilson. Google COVID-19 Community Mobility Reports: Anonymization Process Description (version 1.1). eng (2020).
- [3] M. Barbaro and T. Zeller. A Face Is Exposed for AOL Searcher No. 4417749. *The New York Times* (Aug. 2006). Available from: <https://www.nytimes.com/2006/08/09/technology/09aol.html>.
- [4] D. Blackwell. Equivalent comparisons of experiments. *Ann. Math. Statistics* 24 (1953), pp. 265–272.
- [5] R. Cummings, D. Desfontaines, D. Evans, R. Geambasu, Y. Huang, M. Jagielski, P. Kairouz, G. Kamath, S. Oh, O. Ohrimenko, N. Papernot, R. Rogers, M. Shen, S. Song, W. Su, A. Terzis, A. Thakurta, S. Vassilvitskii, Y.-X. Wang, L. Xiong, S. Yekhanin, D. Yu, H. Zhang, and W. Zhang. Advancing Differential Privacy: Where We Are Now and Future Directions for Real-World Deployment. *Harvard Data Science Review* 6.1 (Jan. 2024).
- [6] Differential Privacy Team, Apple. Learning with Privacy at Scale. Technical Report. Apple, 2017.

-
- [7] B. Ding, J. Kulkarni, and S. Yekhanin. Collecting Telemetry Data Privately. *Advances in Neural Information Processing Systems*. Ed. by I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett. Vol. 30. Curran Associates, Inc., 2017.
 - [8] J. Dong, A. Roth, and W. J. Su. Gaussian differential privacy. *J. R. Stat. Soc. Ser. B. Stat. Methodol.* 84.1 (2022). With discussions and a reply by the authors, pp. 3–54.
 - [9] J. C. Duchi, M. I. Jordan, and M. J. Wainwright. Local Privacy and Statistical Minimax Rates. *2013 IEEE 54th Annual Symposium on Foundations of Computer Science*. 2013, pp. 429–438.
 - [10] C. Dwork. Differential privacy. *Automata, languages and programming. Part II*. Vol. 4052. Lecture Notes in Comput. Sci. Springer, Berlin, 2006, pp. 1–12. ISBN: 978-3-540-35907-4; 3-540-35907-9.
 - [11] C. Dwork, K. Kenthapadi, F. McSherry, I. Mironov, and M. Naor. Our data, ourselves: Privacy via distributed noise generation. eng. *ADVANCES IN CRYPTOLOGY - EUROCRYPT 2006, PROCEEDINGS* 4004 (2006), pp. 486–503.
 - [12] C. Dwork, F. McSherry, K. Nissim, and A. Smith. Calibrating Noise to Sensitivity in Private Data Analysis. eng. *Lecture notes in computer science*. Vol. 3876. Lecture Notes in Computer Science. Berlin, Heidelberg: Springer Berlin Heidelberg, 2006, pp. 265–284. ISBN: 3540327312.
 - [13] Ú. Erlingsson, V. Pihur, and A. Korolova. RAPPOR: Randomized Aggregatable Privacy-Preserving Ordinal Response. eng. *Proceedings of the 2014 ACM SIGSAC Conference on Computer and Communications Security*. New York, NY, USA: ACM, 2014, pp. 1054–1067. ISBN: 9781450329576.
 - [14] P. Kairouz, S. Oh, and P. Viswanath. Extremal mechanisms for local differential privacy. *J. Mach. Learn. Res.* 17 (2016), Paper No. 17, 51.
 - [15] P. Kairouz, S. Oh, and P. Viswanath. The Composition Theorem for Differential Privacy. eng. *IEEE transactions on information theory* 63.6 (2017), pp. 4037–4049.
 - [16] E. Lehmann and J. Romano. Testing statistical hypotheses. eng. 3. ed. Springer texts in statistics. New York, NY [u.a.]: Springer, 2005. ISBN: 9780387988641.
 - [17] A. Narayanan and V. Shmatikov. Robust De-anonymization of Large Sparse Datasets. *2008 IEEE Symposium on Security and Privacy (sp 2008)*. 2008, pp. 111–125.
 - [18] R. Rogers, S. Subramaniam, S. Peng, D. Durfee, S. Lee, S. K. Kancha, S. Sahay, and P. Ahammad. LinkedIn’s Audience Engagements API: A Privacy Preserving Data Analytics System at Scale. eng. *The journal of privacy and confidentiality* 11.3 (2021).
 - [19] L. Steinberger. Efficiency in local differential privacy. 2024.
 - [20] W. J. Su. A statistical viewpoint on differential privacy: hypothesis testing, representation, and Blackwell’s theorem. *Annu. Rev. Stat. Appl.* 12 (2025), pp. 157–175.
 - [21] L. Wasserman and S. Zhou. A statistical framework for differential privacy. *J. Amer. Statist. Assoc.* 105.489 (2010), pp. 375–389.
 - [22] Z. Xu, Y. Zhang, G. Andrew, C. A. Choquette-Choo, P. Kairouz, H. Brendan McMahan, J. Rosenstock, and Y. Zhang. Federated Learning of Gboard Language Models with Differential Privacy. eng. *Proceedings of the conference - Association for Computational Linguistics Meeting*. Vol. 5. 2023, pp. 629–639. ISBN: 9781959429685.