

One-Sided Bounded Noise: Theory, Optimization Algorithms and Applications

Hanshen Xiao
Purdue University / NVIDIA
West Lafayette, IN, USA
hsxiao@purdue.edu

Elaine Shi
Carnegie Mellon University
Pittsburgh, PA, USA
runting@cs.cmu.edu

Jun Wan
Five Rings
New York, NY, USA
junwanthu@gmail.com

Srinivas Devadas
Massachusetts Institute of Technology
Cambridge, MA, USA
devadas@mit.edu

Abstract

We investigate the optimal trade-off between utility and privacy using *one-sided* perturbation. Unlike conventional privacy-preserving statistical releases, randomization for obfuscating side-channel information is often constrained by infrastructure limitations. In practical scenarios, these constraints may only allow *positive* and *bounded* perturbations. For example, extending processing time or sending and storing dummy messages/data is typically feasible. However, implementing modifications in the opposite direction is challenging due to restrictions imposed by hardware capacity, communication protocols, and data management systems. In this paper, we establish the foundation of the positive noise mechanism within three semantic privacy frameworks: Differential Privacy (DP), Maximal Leakage (MaxL), and Probably Approximately Correct (PAC) Privacy. We then present a series of results that characterize or approximate the *optimal* one-sided noise distribution, subject to a second-moment budget and a bounded maximal magnitude. Building on this theoretical foundation, we develop efficient tools to solve the underlying optimization problems. Through experiments conducted in various scenarios, we demonstrate that existing techniques, such as Truncated Biased Laplace noise, are often suboptimal and result in excessive performance degradation. For instance, in an anonymous communication system with a 250K message budget, our optimized DP noise mechanism achieves a $21\times$ reduction in dummy messages and an $18\times$ reduction in dummy message latency overhead compared to traditional methods.¹

CCS Concepts

• Security and privacy;

Keywords

Minimal Constrained Noise, Differential Privacy, Maximal Leakage, PAC privacy, Side-Channel Leakage

¹All the proofs in this paper can be found in the full version https://hanshen-xiao.github.io/files/one_sided_noise.pdf.



This work is licensed under a Creative Commons Attribution 4.0 International License. CCS '25, Taipei, Taiwan

© 2025 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-1525-9/2025/10
<https://doi.org/10.1145/3719027.3765110>

ACM Reference Format:

Hanshen Xiao, Jun Wan, Elaine Shi, and Srinivas Devadas. 2025. One-Sided Bounded Noise: Theory, Optimization Algorithms and Applications. In *Proceedings of the 2025 ACM SIGSAC Conference on Computer and Communications Security (CCS '25)*, October 13–17, 2025, Taipei, Taiwan. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/3719027.3765110>

1 Introduction

Perturbation is one of the most general approaches to randomize a processing procedure for privatizing information disclosure. Various mathematical notions of privacy, including Differential Privacy (DP) [21, 22], Maximal Leakage (MaxL) [27] and PAC Privacy [44] have been established to formally quantify the attainable privacy guarantees through noise addition.

Most commonly-adopted noises are *symmetric*, which, to be more specific, are *unbounded* and *two-sided* centered around 0 (i.e., zero mean). Examples include standard Laplacian noise for pure ϵ -differential privacy (DP) [22], and Gaussian noise for approximate (ϵ, δ) -DP [21] and PAC Privacy [37, 44, 45].

However, many real-world applications impose specific constraints on the noise mechanisms. Due to application semantics, negative noise values may be impermissible. For instance, in anonymous communication systems, it is feasible to inject dummy messages, but removing existing messages could compromise functionality [11, 40]. Similarly, in encrypted database applications, one may pad the running time or the number of memory accesses needed to obscure the query and database contents [14], but reducing them could lead to errors or incorrect outputs. All of the above scenarios can be modeled using *one-sided bounded noise*, i.e., noise that is non-negative, has a positive mean (also referred to as bias), and lies within a bounded interval.

Take DP as an example. The most common practice for adding one-sided noise so far is to rely on a truncated and shifted Laplacian (or geometric) distribution [14, 40]. This approach involves shifting the mean of a standard Laplacian distribution to the positive (right) side until the negative (left) tail becomes sufficiently small, which will be truncated; we then redistribute the probability mass elsewhere after truncation. Unfortunately, this approach does *not* yield optimal error, either for a single disclosure setting or for multiple disclosures that require composition. One key result we will demonstrate is that, the optimal one-sided noise distribution is generally *asymmetric* and heavily *dependent* on the security parameters. This

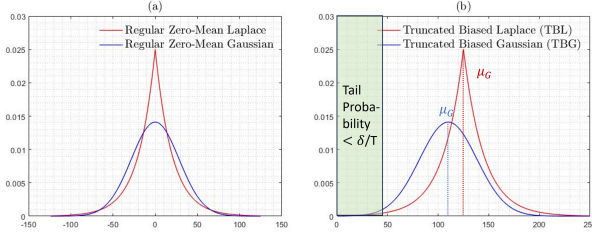


Figure 1: Illustration of zero-mean Laplace/Gaussian Mechanism and positive truncated biased Laplace/Gaussian Mechanism of the same noise variance. (1) Laplace, the taller (red) distribution, is less concentrated with a heavier tail compared to Gaussian, the wider (blue) distribution. (2) To produce the same (ϵ, δ) -DP guarantee, truncated Laplace noise requires a larger bias μ_G but less variance compared to truncated Gaussian noise. In general, the concentration of optimal DP one-sided noise needs to be carefully selected *depending on the given (ϵ, δ) security parameters to balance the bias (the mean μ_G) and variance required, and the optimal noise form is generally not simply either a Laplace or Gaussian.*

is dramatically different from the case of non-biased noise where *independent of* target DP security parameters, the optimal form belongs to or can be closely approximated by some simple distribution class, such as staircase Laplace [26] or Gaussian [5]. Instead, *case-by-case* optimization is necessary for one-sided noise.

The starkly different landscapes are largely due to the different characteristics of the error produced by noise. Throughout this paper, we define the error of injected noise by the expectation of its square (second moment). For one-sided noise, both the bias (i.e., the positive mean) and the variance of the noise distribution contribute to the error, unlike the case of unbiased noise where *only* its variance affects the error. The optimal distribution of one-sided noise thus needs to be carefully selected to balance both bias and variance simultaneously, and, unfortunately, simply shifting classic distributions, including staircase Laplace or Gaussian distributions, is generally *no* longer optimal as partially illustrated in Fig. 1, though they yield (asymptotic) optimal utility (error)-privacy trade-off in the regime of non-biased perturbation, .

In addition to the absence of theory to understand the fundamental gap in utility loss/error caused by the one-sided constraint, to the best of our knowledge, there is also *no* known framework to help privacy practitioners determine the optimal one-sided noise distribution. Note that for two-sided noise, especially in a *compositional* setting with multiple disclosures, many software toolboxes, including OpenDP [25] and Opacus [49], have been developed to automatically optimize noise parameters in standard DP implementations. A library for one-sided noise mechanisms, especially for different privacy metrics, is highly desirable.

1.1 Our Contributions

In this paper, we not only lay the foundation for characterizing the minimal one-sided bounded noise under three semantic privacy definitions—Differential Privacy (DP), Maximal Leakage (MaxL), and PAC Privacy—but also introduce an open-source library, `lib-1sided`

-noise² to assist privacy practitioners in their tasks. In general, the optimal one-sided noise does *not* have a closed-form expression, and the solutions vary across different privacy definitions. To address this, we develop optimization algorithms separately with efficient implementations in our library. These enable automatic determination or approximation of the optimal noise under the selected privacy metric and budget. Additionally, we provide comprehensive comparisons of the semantic interpretations of privacy definitions to guide practitioners in selecting the most appropriate metrics for addressing diverse privacy concerns across applications (see Section 2.4). We summarize our theoretical and algorithmic contributions below:

Differential Privacy For a single release (without composition) with (ϵ, δ) -DP guarantee, we present the closed form of the optimal positive, bounded noise distribution (Theorem 1). For multiple releases, we present a new and tight composition accounting by generalizing the Rényi divergence (Theorem 2) and transforming the determination of minimal positive noise under T -fold composition into a constrained optimization. Finally, we show how to iteratively apply convex optimizations (Algorithm 1) to optimize the noise distribution (Theorem 3) to approximate the optimal solution.

Maximal Leakage We prove that the optimal positive perturbation strategy with minimal cost/overhead to produce $\log(v)$ -MaxL, for an arbitrary real number $v \geq 1$, is a linear interpretation of the optimal deterministic perturbation schemes to produce $\log(\lfloor v \rfloor)$ and $\log(\lceil v \rceil)$ MaxL, respectively (Theorem 4). We then apply dynamic programming to build the *first* efficient algorithm (Algorithm 2) to find the provably optimal perturbation in polynomial time.

PAC Privacy We study the PAC Privacy bound for black-box processing given the output variance. We prove given the second moment budget and a maximal magnitude restriction, the optimal positive noise must be within a family of truncated Gaussian (Theorem 5) and the problem is reduced to optimize the mean and variance of the Gaussian being truncated. We present an efficient optimization algorithm that executes a series of simple iterative binary searches (Algorithm 3).

1.2 Concrete Results on Applications

1.2.1 Network Traffic Leakage. Anonymous communication while hiding who is communicating with whom has received significant attention. It is known that even if messages are encrypted, metadata that reveals active users can be recovered through network traffic analysis [30, 40]. A common privatization mechanism is to ask users to send dummy messages [29], which is essentially a positive perturbation to the packet volume observed by the adversary. Its privacy guarantee has been measured from a DP point of view in [38, 40]. On the other hand, if one adopts dropping messages or partitioning messages into multiple communications to produce a negative obfuscation, it can break soundness unless the system has redundancy built-in, which itself can be expensive.

Compared to truncated Laplace noise adopted in [8, 14, 33, 40], we show the scale of optimized one-sided DP noise can be *orders of*

²<https://github.com/Hanshen-Xiao/lib-1sided-noise>

magnitude smaller in practice: for *Vuvuzela* anonymous communication [40], we achieve $21\times$ less dummy messages with $18\times$ less dummy message latency overhead for a 250K message budget.

1.2.2 Mitigating Cache-Timing Attacks. The execution of cryptographic algorithms or operations on different inputs takes different times and, in practice, it is challenging to write high-performance constant-time software for general-purpose computers [7]. This additional data-dependent timing information could leak cryptographic secrets; in some applications, timing information can be simply characterized by hits and misses to a cache [28]. Bernstein in [7] presented concrete cache-timing attacks to fully recover an Advanced Encryption Standard (AES) secret key. More recently, for OpenSSL's constant-time code, a cache-timing attack on RSA key generation was developed by exploiting a vulnerable code path [2]. Even in Trusted Execution Environments (TEEs), through cache-timing and speculation attacks, Bulck et al. proposed *Foreshadow* to break Intel Software Guard eXtensions (SGX) [39]. In practical execution, it is easy to turn cache hits into misses, but not vice versa. To obfuscate timing and cache information and provably mitigate attacks, the corresponding modification must be positive.

In this paper, we present the first formal analysis with optimal one-sided noise to obfuscate the cache-timing leakage in a 256-bit AES secret key generation from S-boxes [4] under both MaxL and PAC Privacy. In particular, we show a small random positive noise with 1 dummy miss in expectation can provably ensure a negligible adversarial success rate ($< 2^{-245}$) to correctly recover a 256-bit secret key (Fig. 4).

2 Preliminaries

In this section, we formally introduce three semantic and rigorous privacy definitions: Differential Privacy (DP) [20, 22], Maximal Leakage (MaxL) [27], and PAC Privacy [44]. At a high level, the problem of information leakage control can be described by the following generic model: for some sensitive data/input $X \in \mathcal{X}^*$ and some processing function $\mathcal{F} : \mathcal{X}^* \rightarrow \mathcal{Y}^*$, the output $\mathcal{F}(X)$ represents the release/leakage. The goal of privacy preservation is to randomize or modify the original processing function $\mathcal{F}$ into a version $\mathcal{M}$ such that provided the randomized or noisy output $\mathcal{M}(X)$, the adversary cannot implement meaningful inference on the sensitive input X . The following privacy definitions offer formal languages to quantify such hardness from different angles. In Section 2.4, we include a comparison regarding their applicability and the underlying operational challenges in practice.

2.1 Differential Privacy

DEFINITION 1 (DIFFERENTIAL PRIVACY [21]). Given a dataset universe $\mathcal{X}^*$, we say that two datasets $X, X' \subseteq \mathcal{X}^*$ are adjacent, denoted as $X \sim X'$, if X can be obtained by replacing one datapoint in X' , i.e., $X = (X'/x') \cup x$. A randomized algorithm $\mathcal{M}$ is said to be (ϵ, δ) -differentially-private (DP) if for any pair of adjacent datasets X, X' and any event set Y in the output domain of $\mathcal{M}$

$$\mathbb{P}(\mathcal{M}(X) \in Y) \leq e^\epsilon \cdot \mathbb{P}(\mathcal{M}(X') \in Y) + \delta. \quad (1)$$

(ϵ, δ) -DP enjoys an intuitive interpretation where e^ϵ and δ represent a multiplicative and an additive term, respectively, to capture

the worst-case divergence/difference between the likelihood functions produced by two arbitrary adjacent datasets X and X' . From a hypothesis testing perspective, small ϵ and δ will imply either a large Type I or Type II error [19]. In practice, Gaussian and Laplace mechanisms are the workhorses to randomize a processing function for (ϵ, δ) -DP guarantee [23]: the scale of noise is calibrated to the *sensitivity*, i.e., the maximal possible change to the output when one arbitrarily replaces a single datapoint. Unfortunately, tight sensitivity is in general NP-hard to compute [48]. Thus, an important concept in DP research is *composition*, which captures the cumulative privacy risk from multiple releases and plays a key role to privatize algorithms in practice. A complicated algorithm with intractable sensitivity, such as DP-SGD, is usually decomposed into multiple, relatively simpler suboperations with tractable/bounded sensitivity, such as gradient mean estimation of a batch of samples [1, 41, 46, 47]; one can then perturb the intermediate outcomes from each suboperation, assuming that they are released, and derive an upper bound by composing the privacy loss of the release from each iteration, as formalized in the following proposition.

PROPOSITION 1 (Advanced Composition [24]). For any $\epsilon_0 > 0$ and $\delta_0 \in (0, 1)$, the class of (ϵ_0, δ_0) -differentially private mechanisms satisfies $(\epsilon, T\delta_0 + \tilde{\delta})$ -differential privacy under T -fold adaptive composition, where, for any $\tilde{\delta} > 0$,

$$\epsilon = \sqrt{2T \log(1/\tilde{\delta})} \cdot \epsilon_0 + T\epsilon_0(e^{\epsilon_0} - 1). \quad (2)$$

When $\epsilon_0 = o(1/\sqrt{T})$, the latter term $T\epsilon_0(e^{\epsilon_0} - 1)$ is $o(1)$ and the advanced composition (2) roughly states that ϵ scales in $O(\sqrt{T})$ under T -fold composition, which is general and only counts on the (ϵ_0, δ_0) -DP guarantee per release.

2.2 Maximal Leakage

Similar to DP which considers an input-independent guarantee, Maximal Leakage (MaxL) is another operationally-interpretable definition, which measures how much more likely the adversary can identify the true input given the release. Let $U \in \mathcal{U}^*$ denote the user's secret, $X \in \mathcal{X}^*$ denote some intermediate parameter of a processing mechanism $\mathcal{M} : \mathcal{X}^* \rightarrow \mathcal{Y}^*$, whose input and output domains are $\mathcal{X}^*$ and $\mathcal{Y}^*$, respectively, and $Y \in \mathcal{Y}^*$ denote the output of $\mathcal{M}$. Here, $\mathcal{U}^*$, $\mathcal{X}^*$ and $\mathcal{Y}^*$ are all assumed to be finite sets, which can be generalized to the continuous case [27]. Clearly, U, X and Y form a Markov Chain, denoted as $U - X - Y$; the formal definition of MaxL is given below.

DEFINITION 2 (MAXIMAL LEAKAGE [9, 27]). Let $\text{Adv} : \mathcal{Y}^* \rightarrow \mathcal{U}^*$ be an arbitrary algorithm that the adversary applies to recover the secret U from the observation on Y . MaxL with respect to the processing procedure $U - X - Y$ is defined as

$$\mathcal{L}(X \rightarrow Y) = \sup_{U: U-X-Y} \log \frac{\sup_{\text{Adv}} \Pr(U = \text{Adv}(Y))}{\max_u \Pr(U = u)}, \quad (3)$$

which is known to enjoy an equivalent form [3, 27, 32]

$$\mathcal{L}(X \rightarrow Y) = \log \sum_{y \in \mathcal{Y}^*} \max_{x \in \mathcal{X}^*} \mathbb{P}(Y = y | X = x). \quad (4)$$

Compared to DP with a particular focus on *individuals*, the privacy concern of MaxL is regarding the full reconstruction over the *entire* input U , where (3) upper bounds the *multiplicative* gain

to adversary's posterior knowledge on the secret U after observing the release Y . To be more specific, if a mechanism $\mathcal{M}$ satisfies $\log(v)$ -MaxL, then the ratio between the optimal *posterior* chance that an adversary can correctly identify the true input, $\sup_{Adv} \Pr(U = Adv(Y))$, and the optimal *a priori* success rate that the adversary can identify the true input, $\sup_{\tilde{U}} \Pr(U = \tilde{U})$, is bounded by e^v for *any* possible prior distribution of U .

Additionally, the non-adaptive composition of MaxL enjoys a simple summation form [27]. In the same setup $U - X - \tilde{Y} = (Y_1, Y_2)$, let Y_1 and Y_2 be the releases from two mechanisms which are independent conditional on X , then we have

$$\mathcal{L}(X \rightarrow \tilde{Y}) \leq \mathcal{L}(X \rightarrow Y_1) + \mathcal{L}(X \rightarrow Y_2). \quad (5)$$

2.3 PAC Privacy

From an input-independent perspective, Differential Privacy (DP) and Maximal Leakage (MaxL) measures privacy risk/loss in terms of the worst-case distinguishability and the multiplicative gain of successfully identifying the secret, respectively. However, both DP and MaxL require *white-box* algorithmic analysis—such as sensitivity or output likelihood—of the processing function $\mathcal{F}$ to derive provable privacy solutions. This reliance makes it challenging to handle (side-channel) leakage from complicated circuits/processes or involved protocols which generally lack closed form; meanwhile, membership and identification may also not tightly capture broader inference problems, such as (partial) reconstruction attacks.

In contrast, PAC Privacy supports *black-box* privatization. A noise solution can be automatically determined to resist *any* target adversarial inference, and PAC Privacy analysis only requires black-box access to the underlying secret generation and processing function $\mathcal{F}$ [44]. Conceptually, PAC Privacy challenges an adversary to return a *satisfactory* estimation $\tilde{X}$ of the sensitive input X and quantifies the posterior success probability. The criterion of satisfactory reconstruction reflects the level of leakage deemed unacceptable by the secret holder, where, for example, $\tilde{X}$ approximates the salary attribute of a record X with error smaller than 1000, or predicts at least 200 bits correctly of a 256-bit secret key X . The formal definition is given below.

DEFINITION 3 ((δ_ρ, ρ, D) PAC PRIVACY [37, 43–45]). *For a processing mechanism $\mathcal{M} : \mathcal{X}^* \rightarrow \mathcal{Y}^*$, data distribution D , and an inference criterion function $\rho(\cdot, \cdot)$, we say $\mathcal{M}$ satisfies (δ_ρ, ρ, D) -PAC Privacy if the following experiment is impossible:*

A user generates data X from distribution D and sends $\mathcal{M}(X)$ to an adversary. The adversary who knows D and $\mathcal{M}$ is asked to return an estimation $\tilde{X}$ on X such that with probability at least $(1 - \delta_\rho)$, $\rho(\tilde{X}, X) = 1$.

In Definition 3, the probability is based on the randomness in both secret generation $X \leftarrow D$ and the mechanism $\mathcal{M}$. The criterion of satisfactory estimation is captured by the indicator function ρ where $\rho(\tilde{X}, X) = 1$ if and only if $\tilde{X}$ successfully recovers the target feature of X . Continuing with the previous examples, we can define $\rho(\tilde{X}, X) = 1$ iff the estimation error of $\tilde{X}$ on the salary attribute smaller than 1000 or if X and $\tilde{X}$ collide in least 200 bits.

Operationally, to upper bound $(1 - \delta_\rho)$, [44] studies the difference between the *optimal* prior and posterior success rate in f-divergence [36]. Let $(1 - \delta_{o,\rho})$ denote the optimal *a priori* success rate, i.e., the

best chance that an adversary can return a satisfied estimation $\tilde{X}$ such that $\rho(\tilde{X}, X) = 1$ *before* observing the release $\mathcal{M}(X)$:

$$1 - \delta_{o,\rho} = \arg_{X' \in \mathcal{X}^*} \Pr_{X \leftarrow D} (\rho(X', X) = 1). \quad (6)$$

[44] introduces two Bernoulli variables 1_{δ_ρ} and $1_{\delta_{o,\rho}}$, where $\Pr(1_{\delta_\rho} = 1) = (1 - \delta_\rho)$ and $\Pr(1_{\delta_{o,\rho}} = 1) = (1 - \delta_{o,\rho})$, respectively, and considers the f-divergence between these two Bernoulli distributions,

$$\Delta_f^\rho = \mathcal{D}_f(1_{\delta_\rho} \| 1_{\delta_{o,\rho}}) = \delta_{o,\rho} f\left(\frac{\delta_\rho}{\delta_{o,\rho}}\right) + (1 - \delta_{o,\rho}) f\left(\frac{1 - \delta_\rho}{1 - \delta_{o,\rho}}\right), \quad (7)$$

where $f(\cdot)$ can be an arbitrary convex function. It is noted that given the data distribution D and the inference task ρ of interest, the optimal prior rate $(1 - \delta_{o,\rho})$ is determined. Thus, an upper bound of Δ_f^ρ in (7) combined with a lower bound of the prior rate $(1 - \delta_{o,\rho})$ will lead to an upper bound of the target posterior success rate $(1 - \delta_\rho)$. In particular, as a special case of Theorem 1 in [44], when we select the f-divergence to be the KL-divergence, i.e., by selecting $f(t) = t \log t$ in (7), Δ_{KL}^ρ is shown to be bounded by the well-known mutual information [16], as described below.

PROPOSITION 2 ([44]). *For arbitrary ρ and input distribution D ,*

$$\Delta_{KL}^\rho = \delta_\rho \log\left(\frac{\delta_\rho}{\delta_{o,\rho}}\right) + (1 - \delta_\rho) \log\left(\frac{1 - \delta_\rho}{1 - \delta_{o,\rho}}\right) \leq \text{MI}(X; \mathcal{M}(X)), \quad (8)$$

where $\text{MI}(X; \mathcal{M}(X))$ is the mutual information between X and $\mathcal{M}(X)$.

Stemmed from (8), [44] shows how to automatically determine the minimal perturbation ϵ for a black-box processing function $\mathcal{F}$ based on high-confidence estimation on the (co-)variance of $\mathcal{F}(X)$, which ensures that its noisy version $\mathcal{M}(X) = \mathcal{F}(X) + \epsilon$ satisfies provable PAC Privacy guarantees.

2.4 Proper Selection of Privacy Metric

In the previous subsections, we formally introduced three privacy definitions. In practice, before selecting a privacy metric, one should first determine the privacy concern of interest – which part of input data X we aim to protect. The objective could vary from an attribute, to a data point, to relationship between datapoints. DP puts a particular focus on individual privacy and can provide meaningful guarantees especially when the release is an aggregation of multiple individuals, for example, the network traffic when a set of users communicate [40] or the memory pattern when one accesses a database [14], and the goal here is to prevent the adversary from inferring who is talking or which file is accessed. As a comparison, MaxL is not restricted to only make an individual indistinguishable, but also to bound the posterior advantage for an adversary to correctly identify the true input. Thus, MaxL can capture the privacy leakage (adversarial reconstruction hardness) with respect to the *entire* input and can be applied to study the leakage from, for example, processing time or power consumption of a specific program [18], where individual or attribute privacy is *not* meaningful or well-defined. Compared to DP and MaxL, PAC Privacy offers the most general framework to probabilistically describe inference hardness in recovering any specified related information regarding the secret X . However, PAC Privacy requires an *a priori* setup (secret entropy), which is different from the input-independent guarantee in DP and MaxL. In the following, we demonstrate the applicability and operational challenges of the three frameworks.

To randomize a processing function $\mathcal{F}$ to satisfy DP guarantees, one needs to first bound sensitivity, the worst-case change ($\sup_{X \sim X'} \|\mathcal{F}(X) - \mathcal{F}(X')\|$ for arbitrary two adjacent datasets $X \sim X'$) of the output when one arbitrarily replaces a datapoint. As mentioned before, tight sensitivity is intractable in many practical applications, and, usually, one needs to introduce some artificial control, such as clipping and decomposition [1], to produce a sensitivity bound. As for MaxL, sensitivity analysis is not necessary, but MaxL still requires the knowledge of likelihoods across all input selections. Different from DP and MaxL, which cannot view the underlying processing $\mathcal{F}$ as a black box, PAC Privacy enables automated privatization for general inference hardness, not only restricted to distinguishability or identification. However, PAC Privacy requires that the secret distribution D is given or one can repeatedly sample from D . Thus, PAC Privacy is more suitable for statistical data processing or secret key protection, where the input has a clear form of entropy. For example, in protecting an l -bit secret key from cache timing attacks [7, 18, 39], the distribution D of a random secret key X is a uniform distribution over $\{0, 1\}^l$ and thus PAC Privacy is easily applicable. However, when our secret X cannot be sampled or does not enjoy tractable entropy, for example, the messages in anonymous communication, input-independent guarantees becomes the only known feasible solution.

3 A Lesson from Biased Noise –Tradeoff between Mean, Variance and Concentration

In this section, we provide some intuition on the following two important questions: a) why existing positive noise constructions could be sub-optimal, and b) how to construct the optimal positive perturbation. We mainly focus on DP positive noise in this section, but the implications of the results are general, which instruct our following study on the optimal one-sided noise.

In prior works on mitigating side-channel leakage with DP guarantees, Truncated Biased Laplace (TBL) noise [6, 8, 15, 34] and its discrete version, Truncated Geometric noise [13, 50] are among the most-commonly used perturbations. We first take the continuous TBL noise as the example; see definition below.

DEFINITION 4 (TRUNCATED BIASED LAPLACE NOISE [6]). Given parameters $\mu_L > 0$, $\lambda_L > 0$ and $R > 0$, a (μ_L, λ_L, R) Truncated Biased Laplace (TBL) truncates a Laplacian distribution

$$\text{Lap}_{\mu_L, \lambda_L}(z) = \frac{1}{2\lambda_L} \exp\left(-\frac{|z - \mu_L|}{\lambda_L}\right) \quad (9)$$

on range $[0, R]$. The resulted noise e has a probability distribution

$$P(e = z) = \frac{1}{Z_{\mu_L, \lambda_L, R}} \exp\left(-\frac{|z - \mu_L|}{\lambda_L}\right) \cdot \mathbf{1}_{0 \leq z \leq R}. \quad (10)$$

$\mathbf{1}_{0 \leq z \leq R}$ is an indicator which equals 1 when $z \in [0, R]$, otherwise 0. $Z_{\mu_L, \lambda_L, R} = \int_0^R \exp\left(-\frac{|z - \mu_L|}{\lambda_L}\right) dz$ is the normalization parameter.

The following lemma describes the (μ, λ) selection of TBL noise such that it can produce an (ϵ, δ) -DP guarantee.

Lemma 1 (Parameter of Positive Laplace Noise). Suppose a processing function $\mathcal{F} : \mathcal{X}^* \rightarrow \mathbb{R}$ such that for an arbitrary adjacent dataset pair $X \sim X'$, $|\mathcal{F}(X) - \mathcal{F}(X')| \leq s$, i.e., the sensitivity of $\mathcal{F}$ is bounded by s . Then, if we select $\lambda_L = s/\epsilon$, $\mu_L \geq s + \frac{s}{\epsilon} \cdot \log \frac{1}{2\delta(1 - e^{-\mu_L \epsilon/s})}$, and $R = 2\mu_L$, such a (μ_L, λ_L, R) -TBL perturbation ensures (ϵ, δ) -DP.

Intuitively, TBL noise can be viewed as that we perform the following modifications to a standard zero-mean Laplace distribution $\text{Lap}_{0, \lambda}$, defined in (9). First, we shift the $\text{Lap}_{0, \lambda}$ uniformly by μ into $\text{Lap}_{\mu, \lambda}$; second, we truncate its support domain from $(-\infty, \infty)$ to $[0, R]$ and normalize the remaining over $[0, R]$. Started from TBL noise, we have several remarks on positive noise mechanisms and the resultant DP guarantees, compared to the regular zero-mean noise mechanism:

- (1) **ϵ -DP is impossible and a failure probability is necessary.** For an arbitrary positive noise e and an arbitrary deterministic processing function, the respective support sets³ of the distributions of $\mathcal{F}(X) + e$ and $\mathcal{F}(X') + e$, for $X \sim X'$, $\mathcal{F}(X) \neq \mathcal{F}(X')$, cannot be identical. There always exists some subset O such that $\Pr(\mathcal{F}(X) + e \in O) > 0$ while $\Pr(\mathcal{F}(X') + e \in O) = 0$. Once the outputs fall within O , the adversary can perfectly distinguish the input between X and X' , and thus an additive failure rate in a positive noise mechanism is necessary. In TBL, to ensure such a failure probability bounded by δ , we need to select a large enough shift/bias μ_L such that the tail probability between $[0, s]$ or $[R - s, R]$ is bounded by δ , where s is the sensitivity bound.
- (2) **Heavier Utility Loss:** Consider the utility loss captured by the second moment of the injected noise e ,

$$\mathbb{E}[e^2] = \underbrace{(\mathbb{E}[e])^2}_{\text{mean}} + \underbrace{\mathbb{E}[(e - \mathbb{E}[e])^2]}_{\text{variance}}. \quad (11)$$

Given the positive requirement of e , $\mathbb{E}[e]$, the mean of e , must be non-zero, and thus compared to the regular zero-mean noise mechanism, in general, we need to pay an additional utility loss determined by the square of the mean, i.e., $(\mathbb{E}[e])^2$, proportional to $(\mu_L)^2$ in the case of TBL noise.

From Lemma 1, we know that TBL noise behaves as a sufficient method to produce (ϵ, δ) -DP. A natural question is whether TBL is optimal. Before we give a complete answer in the next section, we first show a more intuitive negative answer that TBL is *not* optimal in producing a tight utility-privacy tradeoff under composition. When we consider T -fold composition of a noisy mechanisms using TBL, the cumulative failure probability from the tail scales with T by union bound. Thus, to ensure (ϵ, δ) after T compositions, by advanced composition (Proposition 1) with $\tilde{\delta} = \delta/2$, it suffices to ensure (ϵ_0, δ_0) -DP in each round, where

$$\epsilon_0 = O\left(\frac{\epsilon}{\sqrt{2T \log(2/\delta)}}\right), \text{ and } \delta_0 = \frac{\delta}{2T}. \quad (12)$$

Thus, to ensure (ϵ_0, δ_0) -DP by TBL noise, from Lemma 1, we may select $\lambda_L = \frac{s}{\epsilon_0}$ and $\mu_L = O(\lambda_L \cdot \log \frac{T}{2\delta} + s)$. Consequently, the second moment of the constructed TBL noise is

$$\mathbb{E}[e^2] = O((\mu_L)^2 + (\lambda_L)^2) = O\left(\frac{T \log(1/\delta)}{\epsilon^2} \cdot (\log^2(T/\delta) + 1)\right), \quad (13)$$

for constant s . It is worthwhile noting that the mean/bias of the TBL noise e is $O(\log T/\delta)$ times larger than its standard deviation, as underlined in (13). This matches our intuition that given an $O(e^{-z})$ decaying rate of Laplace distribution, only after an $O(\lambda_L \log 1/(\delta/T))$ length distance from its mean, can we ensure its noise tail is small

³Throughout the paper, we use support set to represent the domain for a distribution, over which there is non-zero probability density/mass.

enough in $O(\delta/T)$. As a comparison, we may similarly consider the Truncated Biased Gaussian (TBG) noise, as an analog of TBL.

DEFINITION 5 (TRUNCATED BIASED GAUSSIAN). Given $\mu_G > 0$ and $\lambda_G > 0$, the probability density function of a (μ_G, λ_G, R) Truncated Biased Gaussian (TBG) noise e is defined as

$$P(e = z) = \mathbf{1}_{0 \leq z \leq R} \cdot \frac{1}{Z_{\mu_G, \lambda_G, R}} \exp\left(-\frac{(z - \mu_G)^2}{2\lambda_G^2}\right), \quad (14)$$

where $Z_{\mu_G, \lambda_G, R} = \int_0^R \exp\left(-\frac{(z - \mu_G)^2}{2\lambda_G^2}\right) dz$ is for normalization.

If we adopt TBG noise to produce (ϵ, δ) -DP under T compositions, similarly it suffices to select $\lambda_G = O\left(\frac{\sqrt{T \log(1/\delta)}}{\epsilon}\right)$, $\mu_G = O(\sqrt{\log(T/\delta)\lambda_G})$ and $R = 2\mu_G$. Then,

$$\mathbb{E}[e^2] = O((\mu_0)^2 + (\lambda_0)^2) = O\left(\frac{T \log(1/\delta)}{\epsilon^2} \cdot (\log(T/\delta) + 1)\right). \quad (15)$$

Comparing (13) with (15), we observe the following:

Square of Mean Larger than Variance: In both positive noise mechanisms with either TBL or TBG, from (13) and (15), to ensure a small tail in a scale δ/T , the bias parameter $\mu_L(\mu_G)$ needs to be a polynomial of $(\log(T/\delta))$ times larger than their standard deviation $\sqrt{T \log(1/\delta)}/\epsilon$, required by the regular (zero-mean) Laplace / Gaussian mechanism to produce the same (ϵ, δ) under T -fold composition. Such a gap *cannot* be simply mitigated by a noise of larger variance: the bias $\mu_L(\mu_G)$ scales with the standard deviation, controlled by $\lambda_G(\mu_G)$, and a larger variance only makes the distributional decay slower, requiring an even larger mean/bias.

Concentration vs. Variance: A closer look at the underlined terms in both (13) and (15)—which capture how many times the mean (or bias) exceeds the standard deviation—reveals that the bias of the TBG noise is smaller than that of the TBL noise by a factor of $O(\log(T/\delta))$. This arises from the fact that the Gaussian distribution is more concentrated. Specifically, the tail of the Gaussian decays at a rate of $O(e^{-z^2})$, compared to $O(e^{-z})$ for the Laplace distribution. This faster decay enables the Gaussian mechanism to use a smaller bias while still achieving the desired tail bound. We illustrate this in Fig. 1 (on page 2).

Indeed, if we expand all constants in (13) and (15), we find that the TBG mechanism requires a larger noise variance than the TBL mechanism. This aligns with our intuition: to achieve the same statistical divergence with a faster-decaying noise distribution, a larger variance is necessary to compensate.

The above two observations suggest that the optimal positive noise distribution requires a careful tradeoff between mean (bias), variance, and concentration. To close the utility-loss gap relative to zero-mean noise, we seek a distribution with sufficient concentration to keep the bias small, while still providing a sharp tail bound to maintain a low failure probability (at most δ/T), guarding against distinguishability via differences in support. At the same time, the distribution must not be overly concentrated, which would result in an excessively large variance.

4 Positive Noise for Differential Privacy

In this section, we systematically study the optimal positive noise mechanism for Differential Privacy (DP). We consider a discrete

processing function $\mathcal{F} : \mathcal{X}^* \rightarrow \mathcal{Y}^* \in \mathbb{Z}$, whose sensitivity is 1, i.e., for two arbitrary adjacent datasets $X \sim X' \in \mathcal{X}^*$, $\sup_{X \sim X'} |\mathcal{F}(X) - \mathcal{F}(X')| \leq 1$. Our goal is to determine the optimal positive noise distribution of e over a bounded set $[0 : R] = \{0, 1, \dots, R\}$ such that its second moment $\mathbb{E}[e^2]$ is minimal while $\mathcal{F}(X) + e$ satisfies an (ϵ, δ) -DP guarantee (under T -fold composition).

Remark 1 (General Sensitivity). In the above setup we normalize the sensitivity to be 1, while all our following results on DP can be generalized to arbitrary integer sensitivity s by leveraging the idea of group DP: if our noise mechanism satisfies (ϵ, δ) -DP with sensitivity being 1, then for the case of sensitivity s , its corresponding DP guarantee scales to $(\epsilon/s, O(\delta/s))$ [23].

4.1 Optimal Positive Noise for a Single Release

We begin by examining the optimal positive noise mechanism for (ϵ, δ) -DP in a single release scenario *without* composition ($T = 1$). We will prove the following general conclusion: given an (ϵ, δ) -DP budget, the optimal noise distribution that minimizes the k -th moment for any $k \in \mathbb{Z}^+$ is *identical* and has a *closed-form* expression. Furthermore, as $R \rightarrow \infty$, meaning that even if there is no bounded restriction for the noise e , the optimal positive noise distribution remains inherently bounded. To be formal, we use $P_e = \{p_0, p_1, \dots\}$ to denote the noise distribution where $p_i = \Pr(e = i)$. The k -th moment of e is defined as $\mathbb{E}[e^k] = \sum_{i=0}^{\infty} i^k \cdot p_i$.

Theorem 1 (Optimum for Single Release). Given a processing function $\mathcal{F}$ of sensitivity 1, among all possible distributions of a positive noise e over $[0, +\infty)$ which ensure an (ϵ, δ) -DP guarantee of the noisy version $\mathcal{F}(\cdot) + e$, the following distribution with probability mass function given in (1) below,

$$p_i = \begin{cases} \delta \cdot e^{\epsilon i} & \text{if } i < \omega \\ \delta \cdot c \cdot e^{\epsilon(2\omega - i)} & \text{if } \omega \leq i \leq \omega', \end{cases} \quad (16)$$

is optimal in a sense that it achieves the minimal k -th moment, for any positive integer k . Here, ω' is either $2\omega - 1$ or 2ω , and $c \in [e^{-2\epsilon}, 1]$ is for normalization such that the sum of p_i equals 1. Here, ω is a turning point, defined as

$$\omega = \frac{1}{\epsilon} \cdot \log\left(\frac{2}{e^\epsilon + 1} + \frac{e^\epsilon - 1}{\delta(e^\epsilon + 1)}\right). \quad (17)$$

Theorem 1 shows that the optimal distribution form is *identical*, in a form of (16), regardless of k . We provide some insights on how the parameters ω, ω' and c are determined. For given ω , we consider the following sequence $\{p_i(\omega)\}$ in a staircase:

$$p_i(\omega) = \begin{cases} \delta \cdot e^{\epsilon i} & \text{if } i < \omega \\ \delta \cdot e^{\epsilon(2\omega - i)} & \text{if } \omega \leq i \leq 2\omega. \end{cases} \quad (18)$$

(18) is the ideal noise shape to ensure an ϵ -multiplicative difference given sensitivity 1. To further ensure $p_i(\omega)$ is a valid probability mass, we need to select ω such that $S(\omega) = \sum_{i=0}^{2\omega} p_i(\omega) = 1$. However, $S(\omega) = 1$ may not have integer solutions, and we address this by setting $\omega = \min\{\omega' \in \mathbb{Z}^+ \mid S(\omega') \geq 1\}$, and introduce a scaling parameter $c \in (0, 1]$ to normalize p_i for $i \in [\omega, 2\omega]$. If $c < e^{-2\epsilon}$, the (ϵ, δ) -DP guarantee can no longer be maintained. In such cases, we adjust the upper bound of i from 2ω to $2\omega - 1$, ensuring that $c \geq e^{-2\epsilon}$.

As a summary, Theorem 1 suggests that given a range restriction R , either no noise distribution within $[0, R]$ can ensure the (ϵ, δ) -DP requirement, or the optimal P_e must be supported on $[0, R_0]$ for some $R_0 \leq 2\omega \leq R$, with ω defined in (17).

In Table 1 and 2, we include numerical results on the second moment of the TBL noise [14] and the optimal noise (from Theorem 1) to achieve the same DP guarantees in various setups. Generally speaking, for stronger DP guarantees (smaller ϵ and δ), Theorem 1 brings more significant improvement.

Table 1: Comparison between the expected square (second moment) of Truncated Bounded Laplace (TBL) Noise and the optimal noise (Theorem 1) for various $(\epsilon, \delta = 10^{-4})$ DP guarantees for a single release.

	$\epsilon = 0.5$	$\epsilon = 1$	$\epsilon = 2$	$\epsilon = 4$	$\epsilon = 8$
TBL	265.5	83.0	27.1	9.8	4.3
Thm 1	244.5(-8%)	68.9(-17%)	17.0(-37%)	4.1(-58%)	1.0(-77%)

Table 2: Comparison between the expected square (second moment) of Truncated Bounded Laplace (TBL) Noise and the optimal noise (Theorem 1) for various $(\epsilon, \delta = 10^{-6})$ DP guarantees for a single release.

	$\epsilon = 0.5$	$\epsilon = 1$	$\epsilon = 2$	$\epsilon = 4$	$\epsilon = 8$
TBL	648.1	187.7	56.3	18.3	6.9
Thm 1	608.1(-6%)	171.5(-8%)	37.7(-33%)	9.1(-50%)	1.1(-84%)

4.2 Hybrid Rényi DP (HRDP)

In the previous section, we have studied the optimality of positive noise for (ϵ, δ) -DP in a *single* iteration. The analysis becomes more involved when we need to further consider composition. Although we introduced advanced composition of (ϵ, δ) -DP in Proposition 1 but we need to mention that (2) is not perfectly tight (both in constants and asymptotically, if we do not ignore the logarithm term [19]). This is especially true when we have additional information on the output distribution of $\mathcal{M}(X)$. For example, when the noise is Gaussian [1], it is known that the (ϵ, δ) metric does not fully characterize the statistical difference between $\mathcal{M}(X)$ and $\mathcal{M}(X')$ to produce the tightest composition. To this end, more involved DP definitions or accounting methods are proposed, such as zero-concentrated DP [10] and Rényi DP [31], which are both established based on Rényi divergence or moment functions. We formally define Rényi DP as follows.

DEFINITION 6 (RÉNYI DIFFERENTIAL PRIVACY [31]). A randomized algorithm $\mathcal{M}$ satisfies $(\alpha, \epsilon(\alpha))$ -Rényi DP (RDP), for some $\alpha > 1$, if for any pair of adjacent datasets $X \sim X'$, $D_\alpha(\mathbb{P}_{\mathcal{M}(X)} \| \mathbb{P}_{\mathcal{M}(X')}) \leq \epsilon(\alpha)$. Here, $\mathbb{P}_{\mathcal{M}(X)}$ and $\mathbb{P}_{\mathcal{M}(X')}$ represent the distributions of $\mathcal{M}(X)$ and $\mathcal{M}(X')$, respectively, and

$$D_\alpha(P \| Q) = \frac{1}{\alpha - 1} \log \int q(y) \left(\frac{p(y)}{q(y)} \right)^\alpha dy, \quad (19)$$

represents α -Rényi Divergence between two distributions P and Q whose density functions are p and q , respectively.

When the output domain $\mathcal{Y}^*$ of $\mathcal{M}$ is discrete, one can simply replace the integral in (19) by summation over elements $y \in \mathcal{Y}^*$ to obtain the discrete RDP version. RDP can be used to elegantly handle the composition of privacy leakage and enables a simple conversion to (ϵ, δ) -DP, as characterized below.

PROPOSITION 3 (RDP Composition and Conversion to (ϵ, δ) DP [31]). For any $\alpha > 1$, the class of $(\alpha, \epsilon_0(\alpha))$ -RDP mechanisms satisfies (ϵ, δ) -differential privacy under T -fold adaptive composition for any ϵ and δ such that

$$\epsilon \geq T\epsilon_0(\alpha) - \log(\delta)/(\alpha - 1). \quad (20)$$

However, unfortunately, RDP cannot be directly applied to handle positive noise mechanisms. As demonstrated in Section 3, with one-sided or bounded noise, the support domains of the output distributions produced by two adjacent datasets, $\mathbb{P}_{\mathcal{M}(X)}$ and $\mathbb{P}_{\mathcal{M}(X')}$, cannot be exactly the same: there always exists some y within the support set of $\mathcal{M}(X)$ but beyond the support set of $\mathcal{M}(X')$ such that $\mathbb{P}(\mathcal{M}(X) = y) \neq 0$ while $\mathbb{P}(\mathcal{M}(X') = y) = 0$, which leads to an unbounded ratio $\mathbb{P}(\mathcal{M}(X) = y)/\mathbb{P}(\mathcal{M}(X') = y) = \infty$ for two adjacent datasets X and X' . Thus, the α Rényi divergence

$$\mathbb{E}_{\mathcal{M}(X')} \left(\frac{\mathbb{P}(\mathcal{M}(X))}{\mathbb{P}(\mathcal{M}(X'))} \right)^\alpha = \int_{-\infty}^{\infty} \frac{(\mathbb{P}(\mathcal{M}(X) = y))^\alpha}{(\mathbb{P}(\mathcal{M}(X') = y))^{\alpha-1}} dy \quad (21)$$

is not well-defined and Proposition 3 is not applicable. To this end, we present a generalization by computing the composition under positive noise in a hybrid form. The high-level idea is to measure the likelihood divergences separately in two cases: the part over the common (overlapped) support set is still measured through Rényi divergence while the remainder is controlled by a failure rate.

We consider an arbitrary processing function perturbed by some one-sided noise, denoted by $\mathcal{M} : \mathcal{X}^* \rightarrow \mathcal{Y}^*$. For an input set X , we use $S_d(X) \subset \mathcal{Y}^*$ to denote the subset of all degenerate events:

$$S_d(X) = \{y \mid \mathbb{P}(\mathcal{M}(X) = y) = 0\}. \quad (22)$$

Accordingly, we define *Partial $(\alpha, \mathcal{R}_{\alpha,p}(X, X'))$ -Rényi Divergence (PRD)* for two adjacent datasets $X \sim X'$ as follows,

$$\mathcal{R}_{\alpha,p}(X, X') = \frac{1}{\alpha - 1} \log \int_{y \in \mathcal{Y}^*/S_d(X')} \frac{\mathbb{P}(\mathcal{M}(X) = y)^\alpha}{\mathbb{P}(\mathcal{M}(X') = y)^{\alpha-1}} dy. \quad (23)$$

Comparing (23) and (19), PRD only measures the divergence within the subset $\mathcal{Y}^*/S_d(X')$ where $\mathbb{P}(\mathcal{M}(X) = y) > 0$. In the following, we present Hybrid RDP (HRDP) to capture the cumulative privacy risk accounting for the release both within or outside the degenerate set S_d . We formally define $(\alpha, \epsilon_{\alpha,p}, \delta_p)$ -HRDP as follows.

DEFINITION 7 (HYBRID RDP). A mechanism $\mathcal{M} : \mathcal{X}^* \rightarrow \mathcal{Y}^* \subset \mathbb{R}^d$ satisfies $(\alpha, \epsilon_{\alpha,p}, \delta_p)$ -HRDP if for arbitrary two adjacent datasets X and X' , the degenerate events are bounded as

$$\sup_{X, X'} \Pr(\mathcal{M}(X') \in S_d(X)) \leq \delta_p;$$

and their PRD defined in (23) is also bounded as

$$\sup_{X, X'} \mathcal{R}_{\alpha,p}(X, X') \leq \epsilon_{\alpha,p}.$$

Theorem 2 (HRDP Composition). For T mechanisms $\mathcal{M}_i, i = 1, 2, \dots, T$ where each $\mathcal{M}_i$ satisfies $(\alpha, \epsilon_{\alpha,p}^{(i)}, \delta_p^{(i)})$ -HRDP, the composition of $\mathcal{M}_{[1:T]}$ satisfies (ϵ, δ) -DP such that for any $\delta' > 0$,

$$\epsilon \geq \sum_{i=1}^T \epsilon_{\alpha,p}^{(i)} + \frac{\log(1/\delta')}{\alpha - 1}, \text{ and } \delta \geq \sum_{i=1}^T \delta_p^{(i)} + \delta'. \quad (24)$$

Compared to the regular composition of RDP described in Proposition 3, Theorem 2 demonstrates the following generalization: one can still apply partial Rényi divergence $\mathcal{R}_{\alpha,p}(X, X')$ between $\mathbb{P}_{\mathcal{M}(X)}$ and $\mathbb{P}_{\mathcal{M}(X')}$ over the non-degenerate domain $\mathcal{Y}^*/S_d(X')$, captured by $\epsilon_{\alpha,p}^{(i)}$, to bound ϵ in (24); on the other hand, the probability over the degenerate set $S_d(X')$, captured by $\delta_p^{(i)}$, is simply additive to the global failure rate δ .

For the discrete case, one can simply replace the integral in (23) by a summation and Theorem 2 still holds. As a final remark, the idea in Theorem 2 can be generalized to other composition accounting methods, for example, through the characteristic function [51], where we can similarly analyze the characteristic function over degenerate and non-degenerate domains, respectively.

4.3 Noise Optimization with Composition

With HRDP, now, we have a more powerful tool and a clearer characterization to handle the DP composition of general noise mechanisms. Given the objective global ϵ and δ bound in (24), respectively, a natural idea to determine the optimal noise distribution is to minimize (24) given the budget on δ , second moment B and maximal magnitude restriction R . However, one remaining obstacle here is that we need to specify the degenerate set S_d in the first place: the objective (24) varies with different selections of S_d . To address this, the following lemma provides a clearer picture on the concentration of the optimal positive noise distribution.

Lemma 2 (Contiguous Support Set). To achieve (ϵ, δ) -DP under T -fold composition, the optimal bounded positive noise $e \in [0, R]$ with the minimal second moment must satisfy the following property: $\Pr(e = 0) > 0$ and if there exists some u such that $\Pr(e = u) = 0$, then $\Pr(e \geq u) = 0$.

Lemma 2 states the following fact: the probability mass of optimal positive noise must be consecutively assigned along some interval $[0, R_0]$ for $R_0 \leq R$. To be specific, for a given R_0 , let $\mathbf{P}_{R_0} = \{p_0, p_1, \dots, p_{R_0}\}$ denote a noise distribution e supported over $\{1, 2, \dots, R_0\}$, where $\Pr(e = i) = p_i$. When sensitivity equals 1, in this discrete setup, the output distributions produced by two adjacent datasets X and X' are either identical or differ by a ± 1 shift, and without loss of generality, we assume $S_d(X) = \{p_{R_0}\}$ and $S_d(X') = \{p_0\}$. Therefore, by Lemma 2, the bound of ϵ in (24) can be equivalently expressed as a function $H(R_0, \mathbf{P}_{R_0})$ of R_0 and $\mathbf{P}_{R_0}$:

$$H(R_0, \mathbf{P}_{R_0}) = \frac{1}{\alpha - 1} \max \left\{ T \log \sum_{i=1}^{R_0} \frac{(p_i)^\alpha}{(p_{i-1})^{\alpha-1}} + \log\left(\frac{1}{\delta - T p_0}\right), \right. \\ \left. T \log \sum_{i=1}^{R_0} \frac{(p_{i-1})^\alpha}{(p_i)^{\alpha-1}} + \log\left(\frac{1}{\delta - T p_{R_0}}\right) \right\}.$$

$H(R_0, \mathbf{P}_{R_0})$ captures the worst case of HRDP by taking the maximal of $\epsilon_{\alpha,p}(X, X')$ with p_0 and $\epsilon_{\alpha,p}(X', X)$ with p_{R_0} . Therefore, we transform determining the optimal noise distribution into the following constrained optimization,

$$\min_{R_0} \min_{\mathbf{P}_{R_0}} H(R_0, \mathbf{P}_{R_0}), \quad (25)$$

$$\text{s.t.} \quad \sum_{i=1}^{R_0} p_i = 1, 0 < p_i < 1, i = 1, 2, \dots, R_0, \quad (26)$$

$$0 < R_0 \leq R, \sum_{i=0}^{R_0} i^2 \cdot p_i \leq B^2, p_0 < \frac{\delta}{T}, p_{R_0} < \frac{\delta}{T}. \quad (27)$$

In (26) and (27), we describe the constraints: (26) ensures that $\mathbf{P}_{R_0}$ is a distribution; (27) ensures that e is supported within $[0, R]$ and the second moment of $\mathbf{P}_{R_0}$ is bounded by B^2 , and the probability of degenerate events S_d is bounded by δ/T . However, (25) is not directly solvable: since as the selection of range R_0 varies, both the form of objective $H(R_0, \mathbf{P}_{R_0})$ and the constraints (26) and (27) change, due to a different $\mathbf{P}_{R_0}$ in S_d . To address this, we consider decomposing the original optimization over $H(R_0, \mathbf{P}_{R_0})$ into the optimization over the tail and R_0 , respectively. We first consider the optimization on $H(R_0, \mathbf{P}_{R_0})$ restricted to noise distributions with fixed R_0 , leftmost δ_l and rightmost δ_r tails. We introduce

$$\mathcal{D}(\delta_l, \delta_r, B, R_0) = \{\mathbf{P}_{R_0} : \sum_{i=0}^{R_0} i^2 p_i \leq B^2, p_0 = \delta_l, p_{R_0} = \delta_r\},$$

to capture the set of distributions supported on $[0 : R_0]$ with second moment bound B and fixed $p_0 = \delta_l$ and $p_{R_0} = \delta_r$.

Theorem 3 (Efficiency of Algorithm 1). Given selections of δ_l, δ_r and R_0 , minimization of $H(R_0, \mathbf{P}_{R_0})$ is equivalent to minimizing

$$\max \left\{ \left(\sum_{i=1}^{R_0} \frac{(p_i)^\alpha}{(p_{i-1})^{\alpha-1}} \right), \left(\sum_{i=1}^{R_0} \frac{(p_{i-1})^\alpha}{(p_i)^{\alpha-1}} \right) \right\},$$

which is convex with respect to $\mathbf{P}_{R_0}$. In addition, given R_0 and $\mathbf{P}_{R_0}$, $H(R_0, \mathbf{P}_{R_0})$ is also convex with respect to p_0 and p_{R_0} , respectively.

The above theorem demonstrates the following facts:

- (a) Given δ_l, δ_r and R_0 , minimizing $H(R_0, \mathbf{P}_{R_0})$ for $\mathbf{P}_{R_0} \in \mathcal{D}(\delta_l, \delta_r, B, R_0)$ is a convex optimization over a convex constraint set $\mathcal{D}(\delta_l, \delta_r, B, R_0)$.
- (b) Given R_0 and $\mathbf{P}_{R_0}$, $H(R_0, \mathbf{P}_{R_0})$ is convex w.r.t. the leftmost and rightmost slot p_0 and p_{R_0} .

To efficiently approximate the optimal noise distribution, we consider fixing $R_0 = R$ and utilizing the convexity shown in Theorem 3 to propose a two-layer algorithm to alternatively optimize (δ_l, δ_r) and $\mathbf{P}_R$, as Algorithm 1.

4.4 Experiments and Comparisons

In this subsection, we set out to produce a set of experiments to show the power of both hybrid RDP (HRDP) accounting (Theorem 2) and optimized noise distribution (Algorithm 1). In practice, the second moment budget B and the maximal magnitude R of noise e capture the expected additional overhead and the worst-case redundancy required, respectively. For example, in anonymous communication [38, 40] or database operations [8, 14], B and R correspond to the expected and maximal dummy messages sent or files written. In

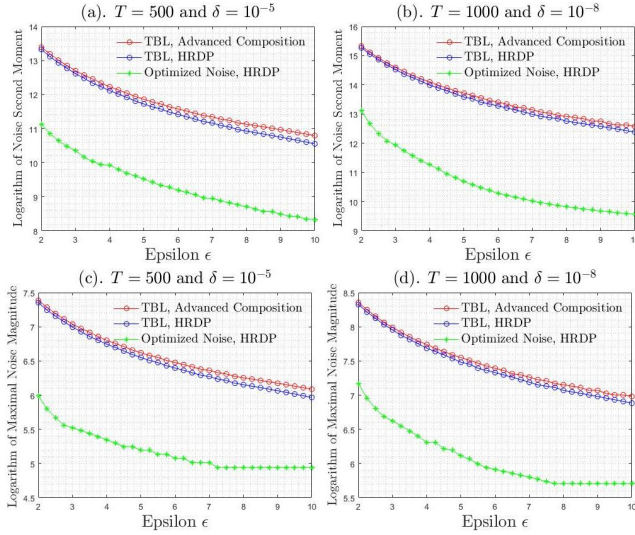


Figure 2: Second moment and maximal magnitude comparison between different positive DP noises and accounting.

Fig. 2(a), under $T = 500$ compositions, we show the required *second moment* $\mathbb{E}[e^2]$ of positive noise to produce (ϵ, δ) -DP with δ fixed to be 10^{-5} in the following three scenarios: 1) by discrete TBL noise in Definition 4 [8, 14, 40, 50] using advanced composition (Proposition 1); 2) still by discrete TBL noise, but using HRDP for composition accounting (Theorem 2); 3) the optimized noise based on HRDP and Algorithm 1. The (1,2,3) cases are captured by the red, blue and green lines, respectively. With a similar setup, in Fig. 2(b), we consider a scenario with more compositions $T = 1,000$ and also a smaller failure rate budget $\delta = 10^{-8}$. Comparing Case 1) (red line) and Case 2) (blue line), we can see HRDP produces a tighter composition bound. Additionally, when we compare Case 2) (blue line) and Case 3) (green line), it can be seen that after optimization, we significantly improve the noise to produce the *same* (ϵ, δ) parameter: the second moment of optimized noise is around $10\times$ and $20\times$ smaller than that of TBL in Fig. 2(a) and (b), respectively. Such improvement is *more significant* with a larger composition T and a smaller privacy budget (ϵ, δ) .

With the same setup as Fig. 2(a,b), in Fig. 2(c,d), we consider the required *maximal magnitude* R for different noises with different accounting methods. Similarly, after optimization, one may find a better noise distribution accommodated in a larger interval with a weaker requirement on R . Compared to TBL noise, the maximal magnitude of optimized noise required is generally $5\times$ smaller in these two examples. To give a more concrete example, we consider the same setup in the anonymous communication protocol *Vuvuzela* [40], where each user has a $T = 250,000$ message budget (composition) for a global budget $\epsilon = \log(2)$, $\delta = 10^{-4}$. By TSL noise and advanced composition [40], the expected and the worst-case number of dummy messages sent per communication round is 300K and 600K, respectively, with an end-to-end additional latency around 18.4 seconds. As a comparison, by selecting a much smaller $R = 30,000$ in Algorithm 1, the optimized noise requires only 14K

Algorithm 1 Optimized Positive Noise for Hybrid RDP

- 1: **Input:** Second moment budget B , range restriction R , composition number T , failure probability budget δ .
- 2: Randomly initialize both tail rate $\delta_l, \delta_r \in (0, \delta/T)$, and accordingly initialize $P_R \in D(\delta_l, \delta_r, B, R)$.
- 3: Alternately run convex optimizer on (δ_l, δ_r) with respect to the loss function $H(R, P_R)$ with updated P_R , and on P_R with respect to the loss function $H(R, P_R)$ given updated (δ_l, δ_r) such that $\delta_l, \delta_r \in (0, \delta/T)$, until they converge.
- 4: **Output:** P_R .

and 30K messages on average and in the worst case, respectively, with latency shortened to 1 second.

5 Positive Noise for Maximal Leakage

5.1 Theory and Algorithm

In this section, we study the optimal positive perturbation for the MaxL measure. We consider a deterministic function $\mathcal{F} : \mathcal{X}^* = \{X_1, X_2, \dots, X_n\} \rightarrow \mathcal{Y}^* = \{1, 2, \dots, m\}$, whose input domain $\mathcal{X}^*$ is formed by m possible selections and the output domain, capturing the information leakage, has n states. A state in $\mathcal{Y}^*$ of larger number (shortened to higher state in the following) has higher overhead, for example, taking a longer processing time, producing more cache misses or requiring a larger memory.

To randomize $\mathcal{F}$ with positive noise, we consider the randomized version of $\mathcal{F}$, denoted by $\mathcal{RF}$, where similarly $\mathcal{RF} : \mathcal{X}^* = \{X_1, X_2, \dots, X_n\} \rightarrow \mathcal{Y}^* = \{1, 2, \dots, m\}$. It is noted that an arbitrary (either deterministic or randomized) processing function can always be represented by a transition matrix $P_{n \times m}$, where the entry at the crossing of the i -th row and j -th column (denoted as $p_{ij} \in [0, 1]$) represents the probability that we map X_i to the j -th state in the output domain $\mathcal{Y}^*$. The positive perturbation constraint requires that the support set of $\mathcal{RF}(X_i)$ is within $[\mathcal{F}(X_i), m]$.

To provide more intuition, we illustrate the corresponding transition matrix $P_{n \times m}$ for the original processing function $\mathcal{F}$ and its positively perturbed version $\mathcal{RF}$, in Fig. 3(a) and Fig. 3(c), 3(d), respectively. We consider an example when $n = m = 7$, and in Fig. 3, each orange box represents a non-zero probability p_{ij} to map X_i to the j -th state. Adding positive noise is intuitively "decomposing" and "moving" the orange boxes in Fig. 3(a) to the right hand side. In Fig. 3(c) and 3(d), we give two positive perturbation schemes. For example, in Fig. 3(c), $\mathcal{F}(X_1) = 1$ (corresponding to $p_{11} = 1$), and for a feasible positive perturbation, $\mathcal{RF}$ is only allowed to map X_1 to higher states, starting from 1 to 7. The perturbed $\mathcal{RF}$ version in Fig. 3(c) changes the mapping of X_1 from state 1 to 2; in Fig. 3(d), the solution becomes that with 0.5 probability $\mathcal{RF}(X_1) = 1$ and with 0.5 probability $\mathcal{RF}(X_1) = 2$.

As defined in (4), the MaxL privacy leakage is measured as $\mathcal{L}(X \rightarrow Y) = \log \sum_{y \in \mathcal{Y}^*} \max_{x \in \mathcal{X}^*} \mathbb{P}(Y = y | X = x)$ which equals $\log(\sum_{j=1}^m \max_i p_{ij})$. Given the transition matrix $P_{n \times m}$, it can be viewed as the logarithm of the sum of the maximal element p_{ij} in each column. In the following, we formalize the cost to produce a mechanism with a satisfied MaxL loss. We introduce an $n \times m$ cost matrix $C_{n \times m}$ where each entry c_{ij} represents the cost of mapping X_i to the j -th output state. Under the positive noise restriction, $c_{ij} = \infty$

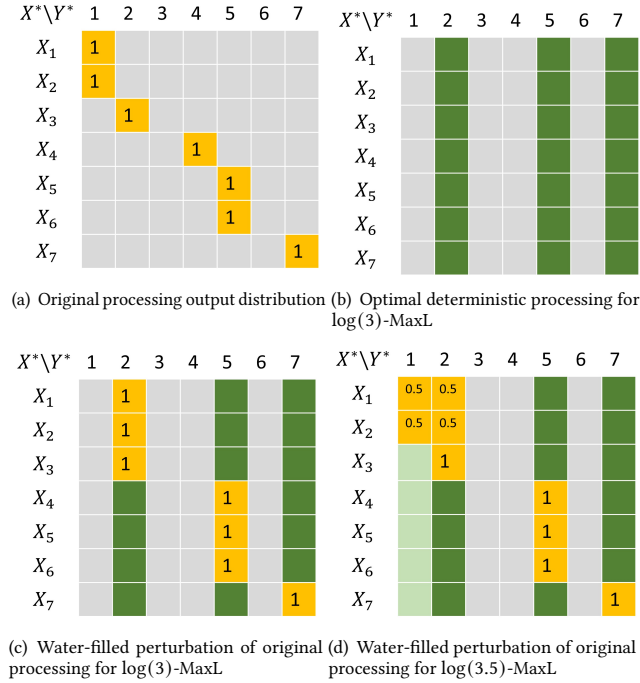


Figure 3: Illustration by transition matrix $P_{n \times m} : \mathcal{X}^* \rightarrow \mathcal{Y}^*$ for positive noise perturbation in MaxL

for $j < \mathcal{F}(X_i)$, i.e., the cost of mapping X_i to lower states compared to the original $\mathcal{F}(X_i)$ is formidably large. In addition, we also assume that for any fixed i , c_{ij} for $j \geq \mathcal{F}(X_i)$ is in a non-decreasing order, i.e., the cost to the higher states is more expensive. Putting the privacy risk and the cost together, determining the optimal perturbation translates to the following constrained optimization problem on transition matrix $P = \{p_{ij}\}_{n \times m}$:

$$\min_P C(P) = \sum_{i=1}^n q_i \sum_{j=1}^m c_{ij} p_{ij} \text{ s.t. } \log \left(\sum_{j=1}^m \max_i p_{ij} \right) \leq \log(v). \quad (28)$$

Here, q_i is the prior distribution of X_i , which captures the frequency that X_i is selected in the processing $\mathcal{F}$.

Before we introduce our main results to determine the optimum of (28), we need to point out a special case of $\mathcal{RF}$ with deterministic modifications to $\mathcal{F}$. In general, to ensure $\log(v)$ -MaxL for an integer $v \geq 1$, a sufficient method is to simply select v states in $\mathcal{Y}^*$ and ensure that the support set of $\mathcal{RF}$ on any X_i is within them. An example is illustrated in Fig. 3(b), where we select $v = 3$ columns (states $\{2, 5, 7\}$) of the transition matrix $P_{n \times m}$ and define an $\mathcal{RF}$ which maps X_i to the *closest* higher state $\{2, 5, 7\}$ compared to $\mathcal{F}(X_i)$ (see Fig. 3(b)). For example, $\mathcal{F}(X_1) = 1$ moves to the closest higher state selected $\mathcal{RF}(X_1) = 2$; $\mathcal{F}(X_4) = 4$ becomes $\mathcal{RF}(X_4) = 5$.

The idea of finding the closest higher state can be further generalized into a *water-filling* operation and the representation of the optimal processing scheme can be further simplified from a transition matrix $P_{n \times m}$ to an m -dimensional vector $\tilde{P}_m$, as studied in [42]. Let $\tilde{P}_m = (\tilde{p}_1, \tilde{p}_2, \dots, \tilde{p}_m)$, where $\tilde{p}_j$ represents the upper bound of p_{ij} in the j -th column. $\tilde{P}_m$ ensures $\log \sum_{j=1}^m \tilde{p}_j$ -MaxL. We say that $\tilde{P}_m$ dominates a transition matrix $P_{n \times m}$ iff $p_{ij} \leq \tilde{p}_j$ for

any i, j . Given a selection of $\tilde{P}_m$, for all transition matrices dominated by $\tilde{P}_m$, the one minimizing the cost function (28) must be in a *water-filling* form, i.e., starting from the state $\mathcal{F}(X_i)$, the optimal mechanism would iteratively fill the j -th slot up to probability $\tilde{p}_j$ for $j = \mathcal{F}(X_i), \mathcal{F}(X_i) + 1, \dots, m$, until the sum $\sum_j p_{ij}$ becomes 1. A formal statement and conclusion is given as follows.

Proposition 4 (Water-filling Lemma [42]). *Among all mechanisms with transition matrix $P = \{p_{ij} | i \in [1 : n]; j \in [1 : m]\}$ dominated by $\tilde{P}_m = (\tilde{p}_1, \tilde{p}_2, \dots, \tilde{p}_m)$, the optimal one with the minimal cost is in the following form: for any i and j , $p_{ij} = \min\{\tilde{p}_j, 1 - \sum_{k=\mathcal{F}(X_i)}^{j-1} \tilde{p}_k\}$.*

Thus, the $\tilde{P}_m$ for a deterministic mechanism must be a binary vector and in particular, if we restrict a feasible positively perturbed $\mathcal{RF}$ to be deterministic for v -MaxL for some integer $v \geq 1$, it is equivalent to selecting v states, denoted by $\mathcal{A} = \{a_1, a_2, \dots, a_v\}$, between $\min_i \mathcal{F}(X_i)$ and $\max_i \mathcal{F}(X_i)$ and sets $\tilde{p}_j = 1$ for $j \in \mathcal{A}$ and it is noted that $\max_i \mathcal{F}(X_i)$ must be selected in $\mathcal{A}$. Therefore, determining the optimal deterministic solution for $\log(v)$ -MaxL with an integer v is reduced to searching over all v -subsets of states with Proposition 4 and comparing their optimal water-filling schemes. A formal definition is given below.

DEFINITION 8 (OPTIMAL DETERMINISTIC SCHEME FOR INTEGER v). *When $v \in \mathbb{Z}^+$, we define DS_v as the set of the deterministic scheme(s) that achieve the minimal cost conditioned on $\log(v)$ -MaxL.*

With the above understanding, there are still two remaining challenges to fully characterize the optimal positive perturbation for MaxL. First and more fundamentally, even when v is integer, is the optimal deterministic scheme also the global optimum for all (possibly randomized) schemes, and for general v , is the optimal solution related to the optimal deterministic schemes in some way? Second, enumeration-based searching takes $O(n \cdot \binom{m}{v})$ time. Can we more efficiently determine the optimal scheme? We will answer both questions affirmatively in the following.

Theorem 4 (Optimal Perturbation for MaxL). *When v is some positive integer, the optimal solution(s) to (28) are exactly DS_v . When $v = \lceil v \rceil - \lambda$ for $\lambda \in (0, 1)$ is not an integer, then the optimal solution(s) to (28) is the linear interpolation of $DS_{\lceil v \rceil}$ and $DS_{\lceil v \rceil + 1}$ as,*

$$\begin{aligned} & \lambda \cdot DS_{\lceil v \rceil} + (1 - \lambda) \cdot DS_{\lceil v \rceil + 1} \\ & = \{\lambda \cdot P_{\lceil v \rceil} + (1 - \lambda) \cdot P_{\lceil v \rceil + 1} \mid P_{\lceil v \rceil} \in DS_{\lceil v \rceil}, P_{\lceil v \rceil + 1} \in DS_{\lceil v \rceil + 1}\}. \end{aligned}$$

Theorem 4 is an improvement of the results in [42] and states the following facts. For $\log(v)$ -MaxL with integer v , the optimal solution should be deterministic and in DS_v . For $\log(v)$ -MaxL with non-integer v , both the cost function and the privacy function are linear with respect to the elements in $DS_{\lceil v \rceil}$ and $DS_{\lceil v \rceil + 1}$, and the optimal solution is in a weighted average of two arbitrary deterministic solutions from $DS_{\lceil v \rceil}$ and $DS_{\lceil v \rceil + 1}$, respectively. Theorem 4 also paves the path to efficiently determine the optimal solution, where it suffices to find optimal deterministic schemes in DS_v (or $DS_{\lceil v \rceil}$ and $DS_{\lceil v \rceil + 1}$). This is solvable by dynamic programming in $O(n \cdot m^2)$ time.

To be specific, we introduce $\mathcal{T}(i, k)$ (sub-algorithm 2 in Algorithm 2) that considers the optimal deterministic mechanism when (1) $\tilde{p}_1, \dots, \tilde{p}_{i-1}$ are given and $\tilde{p}_{i-1} = 1$; and (2) $\sum_{j=i}^m \tilde{p}_j$ equals k . This condition means that given the current selection of $\tilde{p}_1, \dots, \tilde{p}_{i-1}$, we still need to pick k additional states within $[i : m]$. Note that there are totally $n \cdot m$ possible inputs for $\mathcal{T}(\cdot, \cdot)$. For any $\mathcal{T}(i, k)$ such that

$k > 0$, we consider the next state to select in the optimal scheme, i.e., the minimal $j \geq i$ such that $\bar{p}_j = 1$. Once j is given, $\mathcal{T}(i, k)$ is reduced to the sub-problem $\mathcal{T}(j, k - 1)$. This means that $\mathcal{T}(i, k)$ can be solved once we solve $\mathcal{T}(j, k - 1)$ for all $j \geq i$. Therefore, we can use dynamic programming to solve the problem and the time complexity is $O(n \cdot m^2)$.

As an illustration, in Fig. 3(c) and 3(d), we show the optimal positively-perturbed scheme with minimal cost $\min_P C(P)$ for $v = 3$ and $v = 3.5$ when we select the prior distribution $q_i = 1/n$ to be uniform and $c_{ij} = j$ if $j \geq \mathcal{F}(X_i)$, otherwise $c_{ij} = \infty$.

5.2 Applications: AES Secret Key Protection

We provide a concrete application of Algorithm 2 to determine the optimal perturbation scheme to control the MaxL of cache leakage when implementing AES in an S-box implementation [4]. We adopt the same setup and evaluation method as in *Metior* (more details can be found in Section 7 of [18]). Described in MaxL language, we consider $\mathcal{U}^*$ to be a 256-bit secret key space, i.e., $\mathcal{U}^* = \{0, 1\}^{256}$. $\mathcal{X}^*$ corresponds to the intermediate number of distinct cache lines the program touches [12]. Finally, the observation $y \in \mathcal{Y}^*$ captures the number of misses. We use *Metior* [18] to determine both the prior distribution on $\mathcal{X}^*$ and the mapping from $\mathcal{X}^*$ to $\mathcal{Y}^*$. Based on our evaluation in this example, $\mathcal{X}^*$ is formed by 370 intermediate states and $\mathcal{Y}^* = \{1, 2, \dots, 38\}$, i.e., the maximal number of misses is bounded by 38 in all cases. In Fig. 4(a), we consider two kinds of cost functions $C_{n \times m}$. The black line corresponds to the case where $c_{ij} = j$ if $j \geq \mathcal{F}(X_i)$ otherwise $c_{ij} = \infty$, which takes the expected number of misses as the cost. The blue line captures the case where $c_{ij} = j - \mathcal{F}(X_i)$ if $j \geq \mathcal{F}(X_i)$ otherwise $c_{ij} = \infty$, which captures the expected number of *additional* misses. As illustrated by Fig. 4(a), for $\log(v)$ -MaxL, the optimal cost $\min_P C(P)$ is linear when $v \in [\lfloor v \rfloor, \lceil v \rceil]$. Moreover, as v increases, the cost decreases and approaches that of the original processing function. To further interpret the results, it is noted that the prior success rate for an adversary to correctly identify a 256-bit secret key is 2^{-256} . From Fig. 4 we show that at a minimal cost of 6.8 and 1 dummy misses on average, the adversarial posterior success rate to correctly recover the secret key after observing the leakage is bounded, negligibly, by 2^{-253} and 2^{-245} , respectively.

Finally, we want to emphasize that proposed Algorithm 2 is the *first* efficient algorithm to provably determine the optimal scheme in MaxL; prior works either only measure the MaxL of some heuristic protocols without providing privatization solutions [18] or only approximate the optimal perturbation [42].

6 Positive Noise for PAC Privacy

In this section, we study how to determine the optimal positive noise for PAC Privacy. From Proposition 2, we know the posterior advantage measured in KL-divergence $\Delta_{KL}^\rho = \mathcal{D}_{KL}(\mathbf{1}_{\delta_\rho} \| \mathbf{1}_{\delta_{o,\rho}})$, for an arbitrary adversarial inference captured by ρ , is upper bounded by the mutual information $\text{MI}(X; \mathcal{F}(X) + e)$, where X is the sensitive input generated from some distribution D , $\mathcal{F}$ is the processing function and e is some positive noise. As one of the key motivations of PAC Privacy is to enable automatic privatization of black-box processing $\mathcal{F}$, here, we do not put any additional assumptions on the output distribution $Y = \mathcal{F}(X)$ but only assume the variance of

Algorithm 2 Optimal Mechanism for Maximal Leakage

```

1: Input: Objective processing function  $\mathcal{F} : \mathcal{X}^* = \{X_1, X_2, \dots, X_N\} \rightarrow \{1, 2, \dots, m\}$ , prior distribution  $\mathbb{P}_{\mathcal{X}^*}$  of input over  $\mathcal{X}^*$  where  $p_i = \Pr(X = X_i)$ ; cost weight  $c_{ij}$  of mapping  $X_i$  to the state  $j$ ; objective MaxL budget  $\log(v)$ .
2: if  $v$  is integer then
3:   Run Sub-algorithm 1 to determine the optimal deterministic mechanism  $\mathcal{M}_{\mathcal{D}}(v)$  and output  $\mathcal{M}_{\mathcal{D}}(v)$ .
4: else
5:   Run Sub-algorithm 1 to determine the respective optimal deterministic mechanisms  $\mathcal{M}_{\mathcal{D}}(\lfloor v \rfloor)$  and  $\mathcal{M}_{\mathcal{D}}(\lceil v \rceil)$ .
6:   Let  $\lambda = \lceil v \rceil - v$ .
7:   Output  $\lambda \mathcal{M}_{\mathcal{D}}(\lfloor v \rfloor) + (1 - \lambda) \mathcal{M}_{\mathcal{D}}(\lceil v \rceil)$ .
8: end if
```

Sub-algorithm 1: Optimal Deterministic Mechanism $\mathcal{M}_{\mathcal{D}}$. Takes as input an integer $k = v$ and returns the optimal deterministic mechanism in vector form.

```

1: if  $k = 1$  then
2:   Returns  $(0, 0, \dots, 0, 1)$ , which allocates all input to  $m$ .
3: else
4:   Initialize  $i \leftarrow 0$  and  $p \leftarrow (0, 0, \dots, 0)$ .
5:   for  $k'$  in order of  $k, k - 1, \dots, 1$  do
6:      $i \leftarrow \mathcal{T}(i, k').\text{Next}$ .
7:      $p_i \leftarrow 1$ .
8:   end for
9:   Return  $p = (p_1, \dots, p_m)$ .
10: end if
```

Sub-algorithm 2: Dynamic Programming algorithm $\mathcal{T}(a, k)$. $\mathcal{T}$ takes as inputs a position $a \in [m]$ and the remaining budget k .

```

1: if  $k = 1$  then
2:    $\text{Next} \leftarrow m$ .
3:    $\text{Cost} \leftarrow \sum_{(i \mid \mathcal{F}(X_i) \geq a)} p_i \cdot c_{i,m}$ .
4:   Return  $(\text{Next}, \text{Cost})$ .
5: else if  $a \geq m + 1$  then
6:   Return  $(\text{null}, 0)$ .
7: else
8:   for  $a'$  in  $\{a + 1, \dots, m + 1\}$  do
9:      $\text{cost}_{a'} \leftarrow \mathcal{T}(a', k - 1). \text{Cost} + \sum_{a \leq X_i \leq a' - 1} p_i \cdot c_{i,(a' - 1)}$ .
10:   end for
11:    $\text{Next} \leftarrow \arg \min_{a'} \text{cost}_{a'}$ .
12:    $\text{Cost} \leftarrow \min_{a'} \text{cost}_{a'}$ .
13:   Return  $(\text{Next}, \text{Cost})$ .
14: end if
```

$Y, \text{Var}(Y) = \mathbb{E}[(Y - \mathbb{E}[Y])^2]$, is bounded by σ_e^2 . As demonstrated in [44], when the output domain $\mathcal{Y}^*$ of $\mathcal{F}$ is bounded, $\text{Var}(Y)$ can be estimated efficiently in high confidence. Within this general black-box setup, we do *not* assume either the input $\mathcal{X}^*$ or output domain $\mathcal{Y}^*$ of the processing function $\mathcal{F} : \mathcal{X}^* \rightarrow \mathcal{Y}^*$ is finite, and $\mathcal{F}(X)$ for $X \leftarrow D$ can be either discrete or continuous.

In the following, we focus on positive noise e in a continuous distribution within some bounded interval $[0, R]$. Still, the utility loss is defined by the second moment of e , $\mathbb{E}[e^2] = \mu_e^2 + \sigma_e^2$, required to be bounded by some budget B . Here, we use μ_e and σ_e^2 to denote

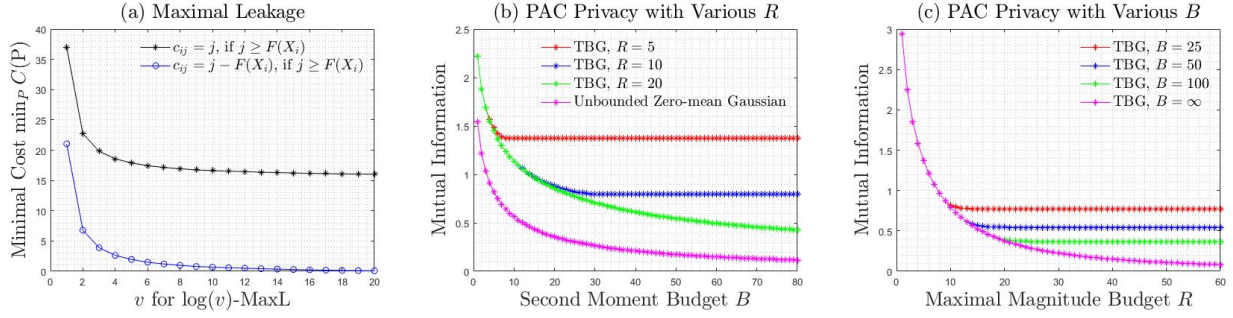


Figure 4: AES Secret Key Protection with Optimal Positive Perturbation in Maximal Leakage and PAC Privacy

Algorithm 3 Generating noise for PAC privacy

- 1: **Input:** The maximum range for the positive noise R ; the second moment budget B ; the standard deviation of output y , σ_y .
- 2: Use binary search to find σ_o that minimizes $\mathcal{L}_{std}(\sigma)$.
- 3: Return $\mathcal{L}_{std}(\sigma_o)$.

Sub-algorithm $\mathcal{L}_{std}$: given a standard deviation σ , find the best mean μ such that the normal distribution $N(\mu, \sigma)$ gives the optimal PAC privacy loss. Return the corresponding PAC loss.

- 1: **if** $S(\frac{R}{2}, \sigma) \leq B$ **then**
- 2: Return $\frac{R}{2}$.
- 3: **else**
- 4: Use binary search to find μ' such that $S(\mu', \sigma) = B$ and return μ' .
- 5: **end if**

the mean (bias) and the variance of the noise e , respectively. To determine the optimal noise distribution D_e , it reduces to solve the following min-max problem,

$$\begin{aligned} & \min_{D_e, \mathbb{E}[e^2] \leq B} \max_{D_{\mathcal{F}(X)}, \text{Var}(\mathcal{F}(X)) \leq \sigma_Y^2} \text{MI}(X; \mathcal{F}(X) + e) \\ &= \min_{D_e, \mathbb{E}[e^2] \leq B} \max_{D_Y, \text{Var}(Y) \leq \sigma_Y^2} h(Y + e) - h(e). \end{aligned} \quad (29)$$

In (29), we use h to represent the differential entropy, where $h(e) = \int_{-\infty}^{\infty} -\mathbb{P}(e = z) \log(\mathbb{P}(e = z)) dz$, and adopt the classic entropy expression of mutual information [16]. (29) models that given a second moment budget B of injected noise e , we aim to minimize the worst-case mutual information for arbitrary processing output distribution $Y = \mathcal{F}(X)$ with bounded variance σ_Y^2 , the only knowledge assumed about the black-box processing function $\mathcal{F}$.

On the other hand, since Y and e are independent, the variance of $Y + e$ is upper bounded by $\sigma_Y^2 + \sigma_e^2$. It is well-known that for any continuous distribution with bounded variance, the Gaussian distribution achieves the maximal entropy, where $h(Y + e) \leq \frac{1}{2} \cdot (\log(2\pi(\sigma_Y^2 + \sigma_e^2)) + 1)$ [16]. Thus, (29) is upper bounded as

$$\min_{D_e, \mathbb{E}[e^2] \leq B} \frac{1}{2} \cdot (\log(2\pi(\sigma_Y^2 + \sigma_e^2)) + 1) - h(e). \quad (30)$$

The following theorem demonstrates that the optimal distribution to (30) under constraints must be a Gaussian distribution $N(\mu, \sigma)$ truncated and normalized within $[0, R]$.

Theorem 5 (Optimal Positive Noise Class for PAC Privacy). *Suppose e is restricted within $[0, R]$ and $\mathbb{E}[e^2] \leq B$, then the optimal noise distribution D_e of (30) must be in a truncated Gaussian form*

$$\mathbb{P}(e = z) = \frac{1}{\sigma(\Phi(\frac{R-\mu}{\sigma}) - \Phi(\frac{-\mu}{\sigma}))} \cdot e^{-(z-\mu)^2/(2\sigma^2)} \cdot \mathbf{1}_{z \in [0, R]},$$

for some $\mu \in [0, R]$, i.e., some Gaussian distribution $N(\mu, \sigma^2)$ conditional on $[0, R]$. Here, $\Phi(t)$ is the cumulative probability function of $N(0, 1)$, i.e., $\Phi(t) = \int_{-\infty}^t 1/\sqrt{2\pi} \cdot e^{-t^2/2} dt$.

By Theorem 5, it suffices to optimize the two parameters μ and σ of a truncated Gaussian distribution within $[0, R]$. Since both σ_e and $h(e)$ in (30) can be expressed using μ and σ , the objective problem with the associated second moment constraint can be rewritten as

In (31), we use the following notations: $\alpha = -\mu/\sigma$ and $\beta = (R - \mu)/\sigma$ with $\varphi(x) = 1/\sqrt{2\pi} \cdot e^{-x^2/2}$ and $\Phi(t) = \int_{-\infty}^t 1/\sqrt{2\pi} \cdot e^{-t^2/2} dt$.

At a first glance, both the objective function $\text{obj}(\mu, \sigma)$ and the second moment constraint $S(\mu, \sigma) \leq B$ in (31) are complicated without a closed form. However, assisted by symbolic analysis in *Mathematica*, we have the following important observations:

a) Given σ , if we ignore the second-moment budget B (or equivalently set $B = \infty$), $\text{Obj}(\mu, \sigma)$ in (31) decreases when $\mu < R/2$ and increases when $\mu > R/2$, and the minimum is achieved when $\mu = R/2$;

b) Given σ , the second moment $S(\mu, \sigma)$ of a Gaussian noise $N(\mu, \sigma^2)$ of mean μ and variance σ^2 truncated over $[0, R]$ strictly increases with μ .

The above two observations suggest that when σ is given, we can adjust the μ value to reduce $\text{Obj}(\mu, \sigma)$ as follows:

- (1) We first evaluate $S(R/2, \sigma)$. If $S(R/2, \sigma) \leq B$, then return $R/2$ as the optimal μ ;
- (2) Otherwise, we need to find the maximum μ such that $S(\mu, \sigma) \leq B$. Given that $S(\mu, \sigma)$ increases with μ , this is equivalent to finding μ such that $S(\mu, \sigma) = B$ and we can use binary search to find an approximation up to any κ accuracy within $\Theta(\log(1/\kappa))$ time.

Now, let $\mu(\sigma)$ denote the optimal μ given σ , and provided a) and b), $\text{Obj}(\mu(\sigma), \sigma)$ also enjoys a nice property, which first decreases and then increases with σ , leading to a unique optimum. This allows a straightforward use of binary search again to find a good σ value, summarized as Algorithm 3.

With Algorithm 3, we continue the study on cache leakage of an AES key in Section 5.2 and present near-optimal positive noise in

$$\min_{\mu, \sigma} \text{Obj}(\mu, \sigma) = \frac{1}{2} \cdot \log(\sigma_Y^2 + \sigma^2) \left[1 - \frac{\beta\varphi(\beta) - \alpha\varphi(\alpha)}{\Phi(\beta) - \Phi(\alpha)} - \left(\frac{\varphi(\beta) - \varphi(\alpha)}{\Phi(\beta) - \Phi(\alpha)} \right)^2 \right] - \log(\sqrt{2\pi e} \cdot \sigma(\Phi(\beta) - \Phi(\alpha))) - \frac{\alpha\varphi(\alpha) - \beta\varphi(\beta)}{2(\Phi(\beta) - \Phi(\alpha))},$$

such that
$$\mathcal{S}(\mu, \sigma) = \left(\mu - \sigma \cdot \frac{\varphi(\beta) - \varphi(\alpha)}{\Phi(\beta) - \Phi(\alpha)} \right)^2 + \sigma^2 \left[1 - \frac{\beta\varphi(\beta) - \alpha\varphi(\alpha)}{\Phi(\beta) - \Phi(\alpha)} - \left(\frac{\varphi(\beta) - \varphi(\alpha)}{\Phi(\beta) - \Phi(\alpha)} \right)^2 \right] \leq B. \quad (31)$$

PAC privacy in Fig. 4 (b,c). In the same setup, following the same evaluation by Metior [18], the variance of released cache misses from a random 256-bit secret key equals $\sigma_Y^2 = 20.6$. In Figure 4(b), we plot the leakage measured in mutual information across various second moment budget B when $R = 5, 10, 20$. Note that given R , the μ of optimal noise must satisfy $\mu \leq \frac{R}{2}$. This is because the selections of (μ, σ) and $(R - \mu, \sigma)$ produce the same privacy loss, but $(R - \mu, \sigma)$ has a larger second moment. Thus, the second moment of the optimal noise is bounded by $R^2/3$ and the privacy loss will decrease with B until $B = R^2/3$. This matches the observation from Figure 4(b), where all curves of TBG noise become flat after some turning points. Meanwhile, there is a gap between the privacy loss produced by positive TBG noise and unbounded, zero-mean Gaussian noise [44], which narrows as R and B increase.

In Fig. 4(c), we plot the privacy loss against the maximum magnitude R , when $B = 25, 50, 100$, and ∞ , respectively. Here, we also observe that the privacy loss stops decreasing with large enough R . Recall that if R is large enough such that $\mathcal{S}(R/2, \sigma) > B$, the optimal μ should be selected such that $\mathcal{S}(\mu, \sigma) = B$. Therefore, after R reaches the threshold determined by $\mathcal{S}(R/2, \sigma) = B$, further increasing R only results in minor changes to the μ and σ of the optimal distribution, and thus provides limited improvement to the privacy risk bound.

7 Additional Related Works

Bounded Noise Mechanism: From an asymptotic viewpoint, the maximal magnitude R of perturbation to produce DP guarantees under compositions has been studied in [17]. In particular, [17] proved that when $\delta \geq e^{-\tilde{O}(T)}$, there exists some bounded noise within $[-R, R]$ such that $R = O(\sqrt{T} \log(1/\delta)/\epsilon)$ to produce (ϵ, δ) -DP under T -fold composition. This suggests that when δ is not too small, the bounded noise mechanism can asymptotically match the same second moment of a general noise mechanism such as Gaussian without constraints. However, though the noise construction of [17] is shown to enjoy the asymptotically-optimal dependence on R , non-asymptotically, its second moment may not be always optimal and can be worse than TSL in the examples considered in Fig. 2. As a comparison, in this paper, we show how to non-asymptotically optimize the noise distribution given various budgets with the more powerful HRDP to handle the composition.

Optimal Mechanism for MaxL: The properties of MaxL and its relationship to other information theory quantities, such as Sibson mutual information, have been studied in [27, 35] and the optimal scheme for MaxL has also been previously studied in [27]. In particular, [27] proved that, with proper non-decreasing assumptions on the cost matrix $C_{n \times m}$, the convex hull of the cost functions determined by deterministic schemes is identical to that formed by both deterministic and randomized schemes. In this paper, we improve their results and show a stronger linear interpolation representation

to fully characterize the optimal scheme for arbitrary $\log(v)$ -MaxL using the optimal deterministic schemes in Theorem 4. Based on Theorem 4, we present the first efficient algorithm to determine the optimal solution in polynomial time.

8 Conclusions and Future Work

In this paper, we studied the optimal positive perturbation in various privacy metrics. With a focus on the utility loss measured by (*weighted*) *second moment* of noise, we provided insights into the characteristics of the optimal perturbation using finite parameters that enable efficient optimization. However, it is worth noting that the expected absolute value of the injected noise may not be the most accurate measure of the resulting utility loss in practice. It would be interesting to generalize our techniques to handle more general utility loss measurements.

We also want to mention, besides mitigating side-channel leakage, one-sided noise is also useful for applications with a specific privacy-preserving overestimation (optimistic) or underestimation (pessimistic) requirement. In addition, for our results on DP, it should be noted that Rényi DP is not the tightest known method to compute composition; even more powerful tools are known, such as, f-DP [19] or through characteristic functions [51]. Thus, one interesting generalization is to consider the hybrid version of those more advanced composition accounting methods. One may follow a similar idea to consider the hybrid versions of those measures on positive noise mechanisms.

For PAC Privacy, in this paper we focused on the general black-box processing without assuming anything specific about the distribution of $\mathcal{F}(X)$ except its variance. When more information is given with respect to $\mathcal{F}$, an interesting problem is to generalize Theorem 5 to accommodate the stronger prior knowledge on $\mathcal{F}$.

Acknowledgments

We thank the anonymous reviewers for their constructive comments.

References

- [1] Martin Abadi, Andy Chu, Ian Goodfellow, H Brendan McMahan, Ilya Mironov, Kunal Talwar, and Li Zhang. 2016. Deep learning with differential privacy. In *Proceedings of the 2016 ACM SIGSAC conference on computer and communications security*. 308–318.
- [2] Alejandro Cabrera Aldaya, Cesar Pereida García, Luis Manuel Alvarez Tapia, and Billy Bob Brumley. 2018. Cache-timing attacks on RSA key generation. *Cryptology ePrint Archive* (2018).
- [3] Mário S Alvim, Konstantinos Chatzikokolakis, Annabelle McIver, Carroll Morgan, Catuscia Palamidessi, and Geoffrey Smith. 2020. *The science of quantitative information flow*. Springer.
- [4] Paulo Barreto Vincent Rijmen and Antoon Bosselaers. [n. d.]. Optimised ANSI C code for the Rijndael cipher. <https://opensource.apple.com/source/BerkeleyDB/BerkeleyDB-15/db/crypto/rijndael/rijndael-alg-fst.c.auto.html>
- [5] Raef Bassily, Adam Smith, and Abhradeep Thakurta. 2014. Private empirical risk minimization: Efficient algorithms and tight error bounds. In *2014 IEEE 55th annual symposium on foundations of computer science*. IEEE, 464–473.

- [6] Johes Bater, Xi He, William Ehrich, Ashwin Machanavajjhala, and Jennie Rogers. 2018. Shrinkwrap: efficient sql query processing in differentially private data federations. *Proceedings of the VLDB Endowment* 12, 3 (2018).
- [7] Daniel J Bernstein. [n. d.]. Cache-timing attacks on AES. ([n. d.]).
- [8] Dmytro Bogatov, Georgios Kellaris, George Kollios, Kobbi Nissim, and Adam O'Neill. 2021. ϵ solute: Efficiently Querying Databases While Providing Differential Privacy. In *Proceedings of the 2021 ACM SIGSAC Conference on Computer and Communications Security*. 2262–2276.
- [9] Christelle Braun, Konstantinos Chatzikokolakis, and Catuscia Palamidessi. 2009. Quantitative notions of leakage for one-try attacks. *Electronic Notes in Theoretical Computer Science* 249 (2009), 75–91.
- [10] Mark Bun and Thomas Steinke. 2016. Concentrated differential privacy: Simplifications, extensions, and lower bounds. In *Theory of Cryptography Conference*. Springer, 635–658.
- [11] Benedikt Bünz, Yuncong Hu, Shin'ichiro Matsuo, and Elaine Shi. 2021. Non-interactive differentially anonymous router. *Cryptology ePrint Archive* (2021).
- [12] Cristian Cadar, Daniel Dunbar, Dawson R Engler, et al. 2008. Klee: Unassisted and automatic generation of high-coverage tests for complex systems programs. In *OSDI*, Vol. 8. 209–224.
- [13] Benjamin M Case, James Honaker, and Mahnush Movahedi. 2021. The Privacy-preserving Padding Problem: Non-negative Mechanisms for Conservative Answers with Differential Privacy. *arXiv preprint arXiv:2110.08177* (2021).
- [14] T-H Hubert Chan, Kai-Min Chung, Bruce Maggs, and Elaine Shi. 2022. Foundations of differentially oblivious algorithms. *ACM Journal of the ACM (JACM)* 69, 4 (2022), 1–49.
- [15] T-H Hubert Chan, Kai-Min Chung, Bruce Maggs, and Elaine Shi. 2022. Foundations of differentially oblivious algorithms. *ACM Journal of the ACM (JACM)* 69, 4 (2022), 1–49.
- [16] Thomas M Cover. 1999. *Elements of information theory*. John Wiley & Sons.
- [17] Yuval Dagan and Gil Kur. 2022. A bounded-noise mechanism for differential privacy. In *Conference on Learning Theory*. PMLR, 625–661.
- [18] Peter W Deutsch, Weon Taek Na, Thomas Bourgeat, Joel S Emer, and Mengjia Yan. 2023. Metior: A Comprehensive Model to Evaluate Obfuscating Side-Channel Defense Schemes. In *Proceedings of the 50th Annual International Symposium on Computer Architecture*. 1–16.
- [19] Jinshuo Dong, Aaron Roth, and Weijie J Su. 2022. Gaussian differential privacy. *Journal of the Royal Statistical Society Series B: Statistical Methodology* 84, 1 (2022), 3–37.
- [20] Cynthia Dwork. 2006. Differential privacy. In *International colloquium on automata, languages, and programming*. Springer, 1–12.
- [21] Cynthia Dwork, Krishnamurthy Korthupadi, Frank McSherry, Ilya Mironov, and Moni Naor. 2006. Our data, ourselves: Privacy via distributed noise generation. In *Advances in Cryptology-EUROCRYPT 2006: 24th Annual International Conference on the Theory and Applications of Cryptographic Techniques*, St. Petersburg, Russia, May 28-June 1, 2006. *Proceedings* 25. Springer, 486–503.
- [22] Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. 2006. 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. Springer, 265–284.
- [23] Cynthia Dwork, Aaron Roth, et al. 2014. The algorithmic foundations of differential privacy. *Foundations and Trends® in Theoretical Computer Science* 9, 3–4 (2014), 211–407.
- [24] Cynthia Dwork, Guy N Rothblum, and Salil Vadhan. 2010. Boosting and differential privacy. In *2010 IEEE 51st Annual Symposium on Foundations of Computer Science*. IEEE, 51–60.
- [25] Marco Gaboardi, Michael Hay, and Salil Vadhan. 2020. A programming framework for opendp. *Manuscript*, May (2020).
- [26] Quan Geng and Pramod Viswanath. 2015. The optimal noise-adding mechanism in differential privacy. *IEEE Transactions on Information Theory* 62, 2 (2015), 925–951.
- [27] Ibrahim Issa, Aaron B Wagner, and Sudeep Kamath. 2019. An operational approach to information leakage. *IEEE Transactions on Information Theory* 66, 3 (2019), 1625–1657.
- [28] Paul C Kocher. 1996. Timing attacks on implementations of Diffie-Hellman, RSA, DSS, and other systems. In *Advances in Cryptology—CRYPTO'96: 16th Annual International Cryptology Conference Santa Barbara, California, USA August 18–22, 1996 Proceedings* 16. Springer, 104–113.
- [29] Albert Kwon, Henry Corrigan-Gibbs, Srinivas Devadas, and Bryan Ford. 2017. Atom: Horizontally scaling strong anonymity. In *Proceedings of the 26th Symposium on Operating Systems Principles*. 406–422.
- [30] David Lazar and Nikolai Zeldovich. 2016. Alpenhorn: Bootstrapping secure communication without leaking metadata. In *12th USENIX Symposium on Operating Systems Design and Implementation (OSDI 16)*. 571–586.
- [31] Ilya Mironov. 2017. Rényi differential privacy. In *2017 IEEE 30th computer security foundations symposium (CSF)*. IEEE, 263–275.
- [32] S Alvim M'rio, Kostas Chatzikokolakis, Catuscia Palamidessi, and Geoffrey Smith. 2012. Measuring information leakage using generalized gain functions. In *2012 IEEE 25th Computer Security Foundations Symposium*. IEEE, 265–279.
- [33] Lianke Qin, Rajesh Jayaram, Elaine Shi, Zhao Song, Danyang Zhuo, and Shumo Chu. 2022. Adore: Differentially Oblivious Relational Database Operators. *Proceedings of the VLDB Endowment* 16, 4 (2022), 842–855.
- [34] Lianke Qin, Rajesh Jayaram, Elaine Shi, Zhao Song, Danyang Zhuo, and Shumo Chu. 2022. Adore: Differentially Oblivious Relational Database Operators. *Proceedings of the VLDB Endowment* 16, 4 (2022), 842–855.
- [35] Sara Saeidian, Giulia Cervia, Tobias J Oechtering, and Mikael Skoglund. 2023. Pointwise maximal leakage. *IEEE Transactions on Information Theory* (2023).
- [36] Igal Sason and Sergio Verdú. 2016. f -divergence Inequalities. *IEEE Transactions on Information Theory* 62, 11 (2016), 5973–6006.
- [37] Mayuri Sridhar, Hanshen Xiao, and Srinivas Devadas. 2025. PAC-Private Algorithms. In *2025 IEEE Symposium on Security and Privacy (SP)*. IEEE.
- [38] Nirvan Tyagi, Yossi Gilad, Derek Leung, Matei Zaharia, and Nikolai Zeldovich. 2017. Stadium: A distributed metadata-private messaging system. In *Proceedings of the 26th Symposium on Operating Systems Principles*. 423–440.
- [39] Jo Van Bulck, Marina Minkin, Ofir Weiss, Daniel Genkin, Baris Kasikci, Frank Piessens, Mark Silberstein, Thomas F Wenisch, Yuval Yarom, and Raoul Strackx. 2018. Foreshadow: Extracting the keys to the intel {SGX} kingdom with transient {Out-of-Order} execution. In *27th USENIX Security Symposium (USENIX Security 18)*. 991–1008.
- [40] Jelle Van Den Hooff, David Lazar, Matei Zaharia, and Nikolai Zeldovich. 2015. Vuvuzela: Scalable private messaging resistant to traffic analysis. In *Proceedings of the 25th Symposium on Operating Systems Principles*. 137–152.
- [41] Di Wang, Hanshen Xiao, Srinivas Devadas, and Jinhui Xu. 2020. On differentially private stochastic convex optimization with heavy-tailed data. In *International Conference on Machine Learning*. PMLR, 10081–10091.
- [42] Benjamin Wu, Aaron B Wagner, and G Edward Suh. 2020. Optimal mechanisms under maximal leakage. In *2020 IEEE Conference on Communications and Network Security (CNS)*. IEEE, 1–6.
- [43] Hanshen Xiao. 2024. *Automated and Provable Privatization for Black-Box Processing*. Ph.D. Dissertation. Massachusetts Institute of Technology.
- [44] Hanshen Xiao and Srinivas Devadas. 2023. PAC Privacy: Automatic Privacy Measurement and Control of Data Processing. In *Annual International Cryptology Conference*. Springer, 611–644.
- [45] Hanshen Xiao, G Edward Suh, and Srinivas Devadas. 2024. Formal Privacy Proof of Data Encoding: The Possibility and Impossibility of Learnable Obfuscation. In *Proceedings of the 2024 ACM SIGSAC Conference on Computer and Communications Security*.
- [46] Hanshen Xiao, Jun Wan, and Srinivas Devadas. 2023. Geometry of sensitivity: twice sampling and hybrid clipping in differential privacy with optimal gaussian noise and application to deep learning. In *Proceedings of the 2023 ACM SIGSAC Conference on Computer and Communications Security*. 2636–2650.
- [47] Hanshen Xiao, Zihang Xiang, Di Wang, and Srinivas Devadas. 2023. A theory to instruct differentially-private learning via clipping bias reduction. In *2023 IEEE Symposium on Security and Privacy (SP)*. IEEE, 2170–2189.
- [48] Xiaokui Xiao and Yufei Tao. 2008. Output perturbation with query relaxation. *Proceedings of the VLDB Endowment* 1, 1 (2008), 857–869.
- [49] Ashkan Yousefpour, Igor Shilov, Alexandre Sablayrolles, Davide Testuggine, Karthik Prasad, Mani Malek, John Nguyen, Sayan Ghosh, Akash Bharadwaj, Jessica Zhao, et al. 2021. Opacus: User-friendly differential privacy library in PyTorch. *arXiv preprint arXiv:2109.12298* (2021).
- [50] Mingxun Zhou, Elaine Shi, T-H Hubert Chan, and Shir Maimon. 2023. A theory of composition for differential obliviousness. In *Annual International Conference on the Theory and Applications of Cryptographic Techniques*. Springer, 3–34.
- [51] Yuqing Zhu, Jinshuo Dong, and Yu-Xiang Wang. 2022. Optimal accounting of differential privacy via characteristic function. In *International Conference on Artificial Intelligence and Statistics*. PMLR, 4782–4817.