



USENIX

THE ADVANCED COMPUTING
SYSTEMS ASSOCIATION

The Adverse Effects of Omitting Records in Differential Privacy: How Sampling and Suppression Degrade the Privacy-Utility Tradeoff

Àlex Miranda-Pascual, *Karlsruhe Institute of Technology and Universitat Politècnica de Catalunya*; Javier Parra-Arnau, *Universitat Politècnica de Catalunya*;
Thorsten Strufe, *Karlsruhe Institute of Technology*

<https://www.usenix.org/conference/usenixsecurity26/presentation/miranda-pascual>

**This paper is included in the Proceedings of the
35th USENIX Security Symposium.**

August 12–14, 2026 • Baltimore, MD, USA

ISBN 978-1-939133-58-8

Open access to the Proceedings of the
35th USENIX Security Symposium
is sponsored by



جامعة الملك عبد الله
للعلوم والتقنية
King Abdullah University of
Science and Technology

The Adverse Effects of Omitting Records in Differential Privacy: How Sampling and Suppression Degrade the Privacy–Utility Tradeoff

Àlex Miranda-Pascual
Karlsruhe Institute of Technology
Universitat Politècnica de Catalunya
alex.pascual@kit.edu

Javier Parra-Arnau
Universitat Politècnica de Catalunya
javier.parra@upc.edu

Thorsten Strufe
Karlsruhe Institute of Technology
thorsten.strufe@kit.edu

Abstract

Sampling is renowned for its privacy amplification in differential privacy (DP), and is often assumed to improve the utility of a DP mechanism by allowing a noise reduction. In this paper, we further show that this last assumption is flawed: When measuring utility at equal privacy levels, sampling as preprocessing consistently yields penalties due to utility loss from omitting records over all canonical DP mechanisms—Laplace, Gaussian, exponential, and report noisy max—, as well as recent applications of sampling, such as clustering.

Extending this analysis, we investigate suppression as a generalized method of choosing, or omitting, records. Developing a theoretical analysis of this technique, we derive privacy bounds for arbitrary suppression strategies under unbounded approximate DP. We find that our tested suppression strategy also fails to improve the privacy–utility tradeoff. Surprisingly, uniform sampling emerges as one of the best suppression methods—despite its still degrading effect. Our results call into question common preprocessing assumptions in DP practice.

1 Introduction

Differential privacy (DP) [15, 17] is firmly established as the state-of-the-art privacy framework, thanks to its strong privacy guarantees and mathematical formulation. A popular technique employed in the DP field is *sampling*, widely regarded for its privacy amplification property [3, 37, 39]: Applying a uniform sampling $\mathcal{S}$ to any (ϵ, δ) -DP mechanism $\mathcal{M}$ yields that the composition $\mathcal{M} \circ \mathcal{S}$ is (ϵ', δ') -DP with $\epsilon' \leq \epsilon$ and $\delta' \leq \delta$. Here, it is commonly assumed that these gains in privacy can be translated into higher utility by calibrating the privacy parameters in $\mathcal{M} \circ \mathcal{S}$ to achieve the same privacy guarantees, while reducing the perturbation applied by the DP mechanism [11, 19, 29]. While the noise introduced by $\mathcal{M}$ can indeed be thus reduced, it remains an open question whether this noise reduction is in general sufficient to offset the utility loss caused by the loss of records through $\mathcal{S}$ itself. In fact,

the contrary has now been demonstrated [35] for DP stochastic gradient descent (DP-SGD) [1]—a well-established and widespread mechanism using sampling.

The question is, beyond DP-SGD, whether the inherent utility loss caused by $\mathcal{S}$ can, in fact, be outweighed by the intended noise reduction. In the composition of $\mathcal{M} \circ \mathcal{S}$, distortion may derive from two principal sources: from the perturbation provided by the protection in $\mathcal{M}$ and from the utility loss caused by the omission of records in $\mathcal{S}$. By translating the privacy amplification into decreased protection demands, we reduce the perturbation of $\mathcal{M}$, thus establishing a tradeoff between these two error sources. In this paper, we investigate the utility effect of sampling in the composition $\mathcal{M} \circ \mathcal{S}$, and importantly, we expand upon these findings to broadly inquire into the possible privacy and utility effects more generally.

To begin, we investigate sampling as a preprocessing step, testing the assumption that it enhances utility through privacy gains. Our experiments compute and compare the utility guarantees of DP mechanisms with and without sampling under identical privacy parameters. We analyze the canonical DP mechanisms (e.g., Laplace, Gaussian, exponential, and report-noisy-max mechanisms [17]) as well as clustering, which has previously been “amplified” by sampling [8]. Our findings reveal that the utility with sampling is worse than that without, indicating that any utility gained from sampling’s privacy amplification does not compensate for the inherent utility loss caused by removing records through sampling in general.

This surprising revelation leads us to ask: *What is the actual effect of deleting records as wanted?* Therefore, we introduce to DP the technique of *suppression* [23], a method from statistical disclosure control (SDC) that targets the deletion of vulnerable records. Suppression is effective in other privacy frameworks, such as k -anonymity [36], where it improves the privacy–utility tradeoff by selectively removing outliers [12, 20]. Similar benefits could then be expected in DP, as outliers also complicate the DP privacy–utility tradeoff.

Therefore, we investigate whether DP mechanisms with suppression yield better privacy or utility than the same mechanism without. Given that outliers may vary in vulnerability

across databases, the decision to delete records can significantly impact the privacy parameters. Therefore, we derive upper bounds on the privacy parameters of $\mathcal{M} \circ \mathcal{S}$ (where $\mathcal{S}$ now denotes suppression) in terms of those of $\mathcal{M}$. We impose no conditions on the suppression algorithm $\mathcal{S}$, thus covering all possible cases, including state-of-the-art sampling.

To assess utility, we replicate our sampling experiments for a family of suppression algorithms, obtaining the same findings. Among the tested mechanisms, our results indicate that DP mechanisms with this suppression do not outperform those without at fixed privacy levels, often yielding worse utility guarantees compared to sampling. Thus, despite the negative outcomes associated with sampling, it remains the superior method for record deletion.

In summary, this paper main contributions are as follows:

- Our experimental study on uniform Poisson sampling over classic unbounded approximate DP mechanisms reveals that, for fixed privacy levels, the utility guarantees of the DP mechanism with sampling are worse than those of the mechanism without sampling.
- We introduce record suppression to DP and we prove how the privacy parameters of $\mathcal{M}$ are affected by preprocessing with any suppression algorithm $\mathcal{S}$, and when we obtain a privacy amplification. To the best of our knowledge, we are the first to provide such a general result for unbounded DP, and additionally, the first to provide a result that is also independent of the choice of $\mathcal{M}$.
- We empirically show that even when factoring in the privacy amplification, our suppression in DP worsens the privacy–utility tradeoff analogously to sampling. We show that, despite both techniques providing unfavorable outcomes, suppression rarely outperforms sampling.

Our findings offer new insights into the relationship between DP, sampling, and suppression. Above all, our findings highlight the need for careful consideration of data preprocessing strategies in privacy-preserving data analysis.

2 Preliminaries and Background

2.1 Differential Privacy

In this paper, we work with (pure) differential privacy (ϵ -DP) [15, 17] and its approximate counterpart, (ϵ, δ) -DP [16, 17]. Their definitions using our notation are as follows:

Definition 2.1 (Differential privacy [17]). Let $\epsilon, \delta \geq 0$ and $\mathbb{D}$ be a class of databases drawn from a data universe $\mathcal{X}$. A randomized mechanism $\mathcal{M}$ with domain $\mathbb{D}$ is (ϵ, δ) -DP if for all neighboring $D, D' \in \mathbb{D}$ and all measurable $A \subseteq \text{Range}(\mathcal{M})$,

$$\mathbb{P}\{\mathcal{M}(D) \in A\} \leq e^\epsilon \mathbb{P}\{\mathcal{M}(D') \in A\} + \delta.$$

If $\delta = 0$, we say that $\mathcal{M}$ is ϵ -DP.

We will work almost exclusively with *unbounded DP*, the original DP notion [15, 17]. Unbounded DP is obtained by selecting in Definition 2.1 the *unbounded neighborhood definition* [25]. Two databases are said to be unbounded neighboring if one is obtained from the other by adding or deleting a record (or, their symmetric difference has size 1: $|D \Delta D'| = 1$). Note that the definition of DP allows for different variants by changing the neighborhood definition [14, 25], such as *bounded DP* [25]. In this case, two databases are said to be *bounded neighboring* if one is obtained from the other by substituting one record for another.

The *privacy parameters*, ϵ and δ , quantify the privacy level of the mechanism $\mathcal{M}$, limiting the amount of information an attacker can extract about the input data. Intuitively, lower values of ϵ and δ provide stronger privacy. Although (ϵ, δ) -DP is defined for all $\epsilon, \delta \geq 0$, only smaller values of ϵ and δ provide reasonable or acceptable privacy levels (consensus would place these bounds at around $\epsilon \leq 2$ and $\delta < \frac{1}{|D|}$ [17]). Any (ϵ, δ) -DP mechanism is also (ϵ', δ') -DP for any $\epsilon' \geq \epsilon$ and $\delta' \geq \delta$. Since larger values provide weaker privacy, it makes sense to find the lowest possible values of ϵ and δ . In particular, given an (ϵ, δ) -DP mechanism $\mathcal{M}$, we say that ϵ and δ are *tight* if there are no $\epsilon'' \leq \epsilon$ and $\delta'' \leq \delta$ (not both equal) such that $\mathcal{M}$ is (ϵ'', δ'') -DP.

2.2 Sampling in DP

Sampling (also known as *subsampling*) is a non-perturbative masking method of SDC that consists in publishing a (random) subset of the original dataset of records [23]. Sampling is also well established for DP, with many theoretical works [3, 29, 37, 39, 41] considering it and studying its effect as a preprocessing algorithm, usually in search of improving the privacy guarantees of DP mechanisms. Balle et al. [3] term this search as the *problem of privacy amplification*: Given a sampling algorithm $\mathcal{S}$ and a DP mechanism $\mathcal{M}$, the goal is to bound the privacy parameters of $\mathcal{M} \circ \mathcal{S}$ by those of $\mathcal{M}$. In particular, $\mathcal{M} \circ \mathcal{S}$ must also be DP, which requires that the sampling technique be well adapted to the neighborhood definition [3]. The expected privacy enhancement is given by the “privacy amplification by sampling” principle [3, 29], which holds that the privacy guarantees of a DP mechanism can be improved, with respect to the original database, when applied to a random subset of records. The rationale behind DP amplification by sampling is as follows: As records are dropped, the privacy should increase, and since there is uncertainty about which records are actually sampled, an attacker will be unable to tell which data has or has not been sampled [39].

DP amplification by sampling was first introduced [37] for *Poisson sampling*, which samples each element x in D with a fixed probability $p \in [0, 1]$. Li et al. [29] later provide the first tight bounds on the privacy parameters of $\mathcal{M} \circ \mathcal{S}$ (proven tight in [3]): If $\mathcal{M}$ is an unbounded (ϵ, δ) -DP mechanism, then $\mathcal{M} \circ \mathcal{S}$ is unbounded $(\ln(1 + p(e^\epsilon - 1)), \delta p)$ -DP (i.e.,

$e^\epsilon - 1$ and δ are reduced by a factor of p). Since the privacy parameters of $\mathcal{M} \circ \mathcal{S}$ are smaller than those of $\mathcal{M}$, the privacy amplification by Poisson sampling is clear. Additionally, Poisson sampling allows, as well, for a slightly more general non-uniform definition, where each element x in a database can be sampled with a different probability $p_x \in [0, 1]$. The tight privacy parameters of $\mathcal{M} \circ \mathcal{S}$ are given by the same formula with $p = \max_{x \in \mathcal{X}} p_x$ [39].

Balle et al. [3] and Steinke [39] provide independent theorems for unbounded and bounded DP. In particular, the theorems provide Poisson sampling for unbounded DP, and *sampling without replacement* (SWOR) [7, 30, 42] for bounded DP, which uniformly samples a subset of D of size $m < |D|$. They prove that, for both sampling algorithms, $\mathcal{M} \circ \mathcal{S}$ is tightly $(\ln(1 + p(e^\epsilon - 1)), \delta p)$ -DP for $p = \max_{x \in \mathcal{X}} \mathbb{P}\{x \in \mathcal{S}(D)\}$. In particular, Balle et al. [3] define their sampling algorithm $\mathcal{S}$ as any algorithm over $\mathbb{D}$ that returns subsets of the input database $D \in \mathbb{D}$. However, to obtain bounds on the privacy parameters of $\mathcal{M} \circ \mathcal{S}$, their theory requires certain assumptions (referred to as “ d_Y -compatibility” in the paper) to be satisfied, which are not achieved for all $\mathcal{S}$ under their definition.

Bun et al. [11] introduce other sampling strategies to pure DP, such as cluster and stratified sampling. The proposed strategies allow more flexibility in sampling, but they all still involve some form of uniform selection. Further, the authors conclude that some sampling strategies cannot enjoy the privacy amplification property. In addition, they present a general theorem that provides a lower bound on the tight parameters of $\mathcal{M} \circ \mathcal{S}$ for most sampling strategies $\mathcal{S}$ in bounded pure DP.

How sampling is applied in DP. Sampling has been widely used to design and improve DP mechanisms. The privacy amplification provided by the most popular variants, Poisson sampling and SWOR, can also be translated into less DP perturbation [11, 19, 29]. Indeed, consider an ϵ -DP mechanism $\mathcal{M}_\epsilon$ (e.g., the Laplace mechanism). By applying Poisson sampling, we can obtain that $\mathcal{M}_\epsilon \circ \mathcal{S}$ is ϵ' -DP with $\epsilon' < \epsilon$, but we can also translate this privacy amplification by reducing the noise (of the Laplace mechanism) as follows: We choose ϵ'' so that $\epsilon = \ln(1 + p(e^{\epsilon''} - 1))$ and replace $\mathcal{M}_\epsilon$ with $\mathcal{M}_{\epsilon''}$ (which adds less noise). The result is $\mathcal{M}_{\epsilon''} \circ \mathcal{S}$ satisfying ϵ -DP for the initial (unmodified) privacy budget ϵ . However, we note here that such noise reduction *does not account for* the potential utility loss caused by sampling records (see Section 3).

In particular, Poisson sampling has been used as a preprocessing step to, for example, “amplify” k -medians and k -means DP clustering mechanisms [8] and dependency-graph generation for the publication of synthetic high-dimensional databases [13]. Furthermore, the literature has also used sampling in alternative settings to preprocessing, like within iterations of mechanisms in machine learning and in the analysis of noisy stochastic gradient descent [1, 24, 26, 33, 34]. For example, DP stochastic gradient descent (DP-SGD) [1] uses multiple iterations of Poisson sampling over the Gaussian mechanism. In the context of stochastic learning with

Rényi DP, sampling has also been studied theoretically for the Laplace mechanism and others [24, 44].

Recently, Räisä et al. [35] studied the effect of sampling on the variance of DP-SGD at fixed privacy levels and concluded that less sampling always leads to a better privacy–utility tradeoff. Even though their conclusion is limited to this concrete utility measure and to binary databases, this calls into question the actual benefit of sampling as a preprocessing step to DP, which we will analyze empirically in the next section.

3 The Effect of Uniform Sampling on Utility

In this section, we study whether the privacy amplification provided by uniform Poisson sampling and translation into less perturbation can indeed provide benefits to the privacy–utility tradeoff of unbounded approximate DP mechanisms.

Thus, in these experiments, we will be comparing the utility levels of a DP mechanism without sampling to those of the same mechanism with sampling—considering the corresponding noise reduction to meet identical privacy guarantees. Our experimentation covers two types of mechanisms: (1) Basic canonical DP mechanisms of DP (i.e., Laplace, Gaussian, exponential, and report noisy max), in the context of statistic computation (precisely, the mean and mode); and (2) Clustering, covering a mechanism [22] that has been previously “amplified” with sampling [8]. For completeness, also note that the Gaussian and Laplace mechanisms have also been studied with sampling [1, 24, 34, 44].

All mechanisms we test allow for rescaling the privacy parameters, essentially calibrating the noise added. Thus, given a DP mechanism and denoting its (ϵ, δ) -DP instantiations as $\mathcal{M}_{\epsilon, \delta}$, we will be comparing the utility values of $\mathcal{M}_{\epsilon, \delta}$ and $\mathcal{M}_{\epsilon'', \delta''} \circ \mathcal{S}$. To ensure that both mechanisms satisfy (ϵ, δ) -DP for the same privacy parameters, it is enough to select $\epsilon'' = \ln(\frac{\epsilon - (1-p)}{p})$ and $\delta'' = \frac{\delta}{p}$ where p is the sampling rate of $\mathcal{S}$. We note that since $\epsilon \leq \epsilon''$ and $\delta \leq \delta''$, the utility loss of $\mathcal{M}_{\epsilon'', \delta''}$ is, generally, equal or lower than that of $\mathcal{M}_{\epsilon, \delta}$; yet, the question remains on how the utility losses of $\mathcal{M}_{\epsilon, \delta}$ and $\mathcal{M}_{\epsilon'', \delta''} \circ \mathcal{S}$ compare. Note that comparing the utility guarantees under the same privacy level allows for a fair comparison of the tradeoff.

3.1 Experiment Setup

Mean computation. We protect the mean in two ways: Using the Laplace and Gaussian mechanisms. We consider two independent $\frac{\epsilon}{2}$ -DP Laplace (or $(\frac{\epsilon}{2}, \frac{\delta}{2})$ -DP Gaussian) mechanisms, one for the sum query f_{sum} that sums all the values in the database, and one for the counting query f_{count} that counts how many records are in the database. The noisy mean is then obtained by dividing the noisy sum by the noisy count, which is an ϵ -DP mechanism (or (ϵ, δ) -DP) by sequential composition and post-processing [17]. This variation, called

NoisyAverage, is a well-known DP mechanism to compute the mean that reduces the overall noise that would be necessary to protect the mean query function directly with Laplace or Gaussian noise [28].

Mode computation. We compute the mode in four ways using report noisy max (RNM) and the exponential mechanism [17]. RNM is used to determine which of the k count queries f_k has the maximum value, and thus we can use it to return a perturbed mode with DP protection. RNM achieves ϵ -DP by adding Laplace noise: For all $D \in \mathbb{D}$, it is defined as $\mathcal{M}_{\text{RNM}}^{f, \epsilon}(D) = \arg \max_{i \in [k]} \{f_i(D) + z_i\}$ with $z_i \sim \text{Lap}(\frac{\Delta f_i}{\epsilon})$ (i.i.d.). A variation of RNM is obtained by adding exponential noise from $\text{Exp}(\frac{\epsilon}{2\Delta f_i})$ instead, which still satisfies ϵ -DP. In our case, the query functions act over disjoint support (i.e., the elements of $\mathcal{X}$), and thus RNM can be viewed as a parallel composition [31] of $|\mathcal{X}|$ Laplace mechanisms. This fact allows us to obtain an (ϵ, δ) -DP variant using Gaussian mechanisms (note that this is not generally true for RNM with Gaussian noise [27]). In addition, we also protect the mode using the exponential mechanism by defining the score function as the count query minus the maximum value in $\mathcal{X}$ (this last term ensures the score function is negative, avoiding computational inaccuracies caused by large floating-point numbers).

Clustering mechanisms. Recall that DP clustering has previously been amplified via sampling [8]. Our experiment covers their tested k -median algorithm [22], achieving DP through the exponential mechanism. We also test a different k -means algorithm (due to some ambiguity of the original method). We choose the well-known DP version of the k -means clustering (i.e., Lloyd’s algorithm) introduced by Blum et al. [9] for this purpose. This mechanism, also known as DPLloyd [40], achieves DP by computing the centroids in each iteration with the Laplace NoisyAverage mechanism over each cluster. We refer to Su et al. [40] for further details.

Utility metrics. To keep our plots consistent, we ensure that for all utility metrics $u(\mathcal{M}, D)$, larger values indicate worse utility (increased errors), and values close to 0 indicate better utility preservation. We are thereby providing intuition on the amount of error or inaccuracy. For the mean computation, we stick to common practice and take $u(\mathcal{M}, D)$ as the mean percent error (MPE) between the real mean $\frac{f_{\text{sum}}(D)}{f_{\text{count}}(D)}$ and the output noisy mean. For the mode computation, we take $u(\mathcal{M}, D)$ as the probability of incorrectly returning the argument of the maximum of D . For the k -median clustering mechanisms, we take the average of the L^2 distances of each record to the closest median (the average *cost* [22]). Finally, for DPLloyd, we take the normalized intracluster variance (NICV), defined as the average of the squares of the L^2 distance of each record to the centroid of the assigned cluster. NICV is a common metric used to evaluate k -means clustering approaches including DPLloyd [40].

Databases. For the computation of mean and mode, we consider three well-known popular numerical databases in

Database	Columns	Value range	Sensitivity bounds
Adult [6] (Size: 32 561)	age	17–90	0–125
	hours-per-week	1–99	0–100
Census [10] (Size: 1 080)	FEDTAX	1–21 260	0–31 889
	FICA	6–7 932	0–11 890
Irish [2] (Size: 66 666)	Age	15–84	0–125
	Education	1–10	1–10

Table 1: Databases employed in the experimentation.

the field of SDC. For each database, we select two columns to use in our evaluations, considering each column as its own one-dimensional database. Table 1 shows the selected databases and columns, where we prioritized different numerical ranges for variability and simpler-to-understand attributes for each database (such as ages). However, we do not compute the mode over the columns of the Census database because multiple elements reach the maximum count for each column (in particular, no element repeats in FEDTAX).

In our computations, we will need bounds on the values of each column (e.g., to compute the sensitivity of the Laplace/Gaussian mechanism). Since DP is a property that does not depend on the choice of database, lower and upper bounds are usually chosen that do not necessarily match the range of values in the database. Following field practices [38], we either select logical extremal bounds (e.g., 0 to 125 for ages) or 0 to $\lceil 1.5 \text{max_value_in_database} \rceil$ if no clear upper bound exists—note that this does not constitute a privacy violation, but the contrary, it is an estimation of the possible domain range meant to represent every database in $\mathbb{D}$ [38]. The exact values chosen are shown in Table 1.

We run DPLloyd on the Adult database [6] under the same conditions as Su et al.’s experiment [40]: The clustering is performed over the six numeric columns of the database and for $k = 5$ clusters. All values are (min-max) normalized to $[-1, 1]$ as required by DPLloyd. The chosen k -median algorithm is not empirically evaluated in the original publications, but only theoretically [8, 22]. Therefore, following the mechanism requirements and due to large computational cost, we first generate a random two-column database over $\{1, \dots, 100\}^2$. We sample 100 points using a Gaussian distribution with $\sigma = 10$ (nearing to the closest integer) centered at four randomly selected accumulation points in $\{10, \dots, 90\}^2$. The database is then normalized so the sensitivity is 1 and we select $k = 4$.

3.2 Experiments and Results

For every database and mechanism, we compute the utility metric values of the mechanism with and without sampling for various privacy parameters and sampling rates. We run the experiments for $\epsilon \in \{0.25, 0.5, 1, 2\}$. We use the optimal Gaussian mechanism [4] that, unlike the classic version [17], is also defined for $\epsilon \geq 1$. In addition, parameter δ is set to

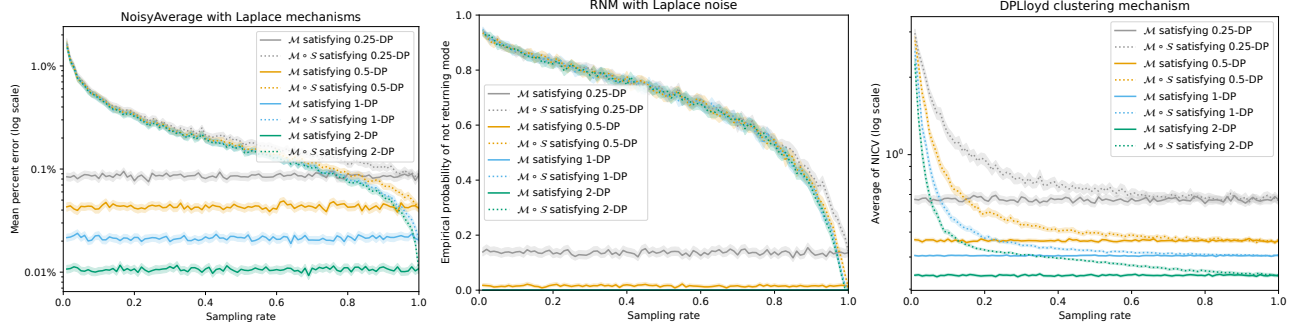


Figure 1: Plots of the utility values of $\mathcal{M}$ and $\mathcal{M} \circ \mathcal{S}$ for the uniform Poisson sampling for the Adult database (and age column). The shaded areas correspond to a 95% confidence interval (CI) for the mean of the utility metric (95% Wilson CI for the mode).

$|D|^{-2}$ when working with the Gaussian mechanism (as suggested in the literature [17]), the only mechanism that requires a non-zero δ . We test every sampling algorithm $\mathcal{S}_p$ with sampling rate $p \in \{0.01, 0.02, \dots, 0.98, 0.99\}$. Since all mechanisms and sampling algorithms are randomized, for each instantiation, we compute $u(\mathcal{M}_{\epsilon, \delta}, D)$ and $u(\mathcal{M}_{\epsilon'', \delta''} \circ \mathcal{S}_p, D)$ 500 times (mean and DPLloyd), 2 000 times (mode), or 20 times (k -median), and we always provide their means.

Our results are surprising as they show that the utility guarantees of $\mathcal{M}_{\epsilon'', \delta''} \circ \mathcal{S}_p$ are worse than those of $\mathcal{M}_{\epsilon, \delta}$ for all mechanisms, databases, privacy parameters ϵ and δ , and (almost all) sampling rates p we tested. Räisä et al. [35] make a similar observation for DP-SGD, but given the general interest in privacy amplification [11, 19, 29], we consider it rather surprising that our results hold for all tested mechanisms.

In Figure 1¹, we show that the utility values differ significantly for most sampling rates, with $\mathcal{M} \circ \mathcal{S}$ having worse utility than $\mathcal{M}$ (here, simplifying the notation). The difference becomes small near a sampling rate of 1, and we find a few rates where $\mathcal{M} \circ \mathcal{S}$ preserves utility better than $\mathcal{M}$. We assume that this is due either to issues of floating-point precision, or rare beneficial random choices by the sampling algorithm.

Looking at the mechanisms independently, we note that the mean computation with sampling provides very small error, remaining less than 0.25% MPE for sampling rates larger than 0.4 and less than 2% even for more abrasive rates near 0; yet, it always increases with respect to the mechanism without sampling. The mode computation expresses variations across the databases depending on the original distribution. For example, in the `hours-per-week` column in the Adult database, the mode represents more than half the results in the database, and thus both $\mathcal{M}$ and $\mathcal{M} \circ \mathcal{S}$ provide a perfect failure probability of 0. For more varied data, like the `age` column in the Adult database, we see drastic utility losses with sampling, increasing the failure probability of RNM (Laplace) from under 18% to over 60% for most sampling rates. DPLloyd also exhibits a utility degradation under sampling, but contrary

to the others, $u(\mathcal{M} \circ \mathcal{S}, D)$ remains quite close to $u(\mathcal{M}, D)$ until spiking at around $p = 0.1$. The k -median clustering also shows a utility degradation similar to the previous plots; however, this does not contradict the theoretical utility evaluations performed on sampling [8] in which the sum of distances (rather than the average sum) is compared, resulting in bias with respect to the database and sample size.

In summary, our result shows that, for the tested mechanisms, it is preferable to apply mechanism $\mathcal{M}$ directly under the target privacy parameters than to rely on the privacy amplification of sampling to improve the privacy–utility tradeoff.

4 Introducing Suppression to DP

Our previous experimental results reveal that uniform Poisson sampling has detrimental effects on the utility that DP mechanisms yield: The utility gain from translating the privacy parameters is insufficient to counteract the utility loss caused by omitting records.

Yet, records in databases are complex and diverse, and they have different degrees of vulnerability. Furthermore, experience shows that some records are harder to protect than others or have different costs for the protecting mechanism. Thus, any loss in utility that is caused by omitting records could be reduced when records are omitted strategically, for example, by targeting the hardest-to-protect or outlying records. This process is known as *suppression* [23], and it has been shown as a mechanism amplifier under syntactic privacy notions like k -anonymity [18, 20]. Nevertheless, to the best of our knowledge, no studies treat the effect of suppression for DP.

Thus, in the following, we introduce suppression to DP. Formally, suppression is a SDC non-perturbative masking technique like sampling [23]. In suppression, data values are deleted from the original dataset to eliminate easily identifiable features. Our suppression corresponds to *whole-record suppression* [5, 18, 20, 36, 43], but we note that suppression can also refer to deleting specific data values of the records [23].

We will formalize DP suppression as a generalization of sampling. While the state of the art on sampling in unbounded

¹We provide all plots in the long version ([arXiv:2601.05180](https://arxiv.org/abs/2601.05180)).

DP works exclusively with algorithms defined according to a uniform selection scheme [18], we define suppression completely general, covering any way of choosing or omitting records. Nevertheless, we acknowledge that our definition of suppression in DP does match the general definition of sampling by Balle et al. [3]. We will provide bounds on the privacy parameters for all suppression algorithms and empirically evaluate a family of suppression algorithms. In this way, we address an open question in the literature: What are the effects of sampling/suppression when defined more flexibly (e.g., in a non-uniform manner) on DP mechanisms?

4.1 The Suppression Algorithm

We now turn to investigating the effect of deleting some records on a DP mechanism. To keep our results well-defined, we assume that $\mathbb{D}$ is closed by subsets (or subdatabases), i.e., if $D \in \mathbb{D}$ and $C \subseteq D$, then $C \in \mathbb{D}$. This is the common assumption in the literature [17]. Suppression is modeled by a *suppression algorithm* $\mathcal{S}$ with domain $\mathbb{D}$ that, given any database $D \in \mathbb{D}$, deterministically or randomly outputs a subset of D or, equivalently, suppresses a subset of D . Since databases $D \in \mathbb{D}$ are of finite size, we consider $\mathcal{S}(D)$ to be a discrete random variable that outputs subsets of D (including the case where $\mathcal{S}$ is a deterministic function). Note that our definition does not impose any restrictions on the deletion process, nor does it establish relations between $\mathcal{S}(D_1)$ and $\mathcal{S}(D_2)$ for different $D_1, D_2 \in \mathbb{D}$. In particular, completely different suppression techniques can be defined independently for each database in $\mathbb{D}$. Thus, $\mathcal{S}$ is defined in a completely general way without any additional restrictions.

We introduce two types of suppression: database-dependent and database-independent. In *database-independent* suppression, an element or subset of elements is deleted with the same probability regardless of the database to which it belongs, i.e., $P\{C \subseteq \mathcal{S}(D_1)\} = P\{C \subseteq \mathcal{S}(D_2)\}$ for all $D_1, D_2 \in \mathbb{D}$ and for all $C \subseteq D_1, D_2$. Some examples of database-independent suppression consist of deleting records—deterministically or probabilistically—over or under a predefined threshold (e.g., deleting all individuals over the age of 100 or over the height of 2.10 m). In particular, any same record is deleted with the same probability across all different databases (i.e., $P\{x \in \mathcal{S}(D_1)\} = P\{x \in \mathcal{S}(D_2)\}$ for all $x \in D_1, D_2$). We refer to all other types of suppression as *database-dependent*, where a same record may be deleted with different probabilities in two databases. This includes deleting records according to their mean or counts in the database, which varies across databases.

We note that classic state-of-the-art sampling algorithms [3, 29, 39] are all database-independent, with many [3, 29] also being *uniform* (i.e., $P\{x \in \mathcal{S}(D)\}$ is equal for all $D \in \mathbb{D}$ and $x \in D$). Bun et al. [11] introduce specific database-dependent sampling algorithms, but these contain some kind of uniform selection and are thus more limited in our suppression context.

4.2 The Suppression Problem

We extend the *privacy amplification problem* of sampling [3] in order to define the *suppression problem* as follows: Given an unbounded (ϵ, δ) -DP mechanism $\mathcal{M}$ and a suppression algorithm $\mathcal{S}$, both with domain $\mathbb{D}$, what are the privacy and utility guarantees provided by $\mathcal{M} \circ \mathcal{S}$? In particular, with this problem, we are interested in understanding under which conditions $\mathcal{M} \circ \mathcal{S}$ has beneficial properties, and whether they improve over those of $\mathcal{M}$. Posing these questions, we are interested in understanding exactly when $\mathcal{M} \circ \mathcal{S}$ also satisfies approximate DP, and with which privacy parameters. In addition, given our experimental results (see Section 3), we highlight the importance of knowing the effect on the output utility of $\mathcal{M} \circ \mathcal{S}$ compared to that of $\mathcal{M}$.

The suppression problem already becomes relevant when discussing database-independent vs. -dependent suppression. Database-dependent suppression is better suited to the deletion of vulnerable or outlier records because these records usually depend on the database to which they belong. However, such records impose a cost on the privacy parameters. In essence, DP must protect or take into account any change between databases, which consumes privacy budget. For these reasons, differences between databases which, on the one hand, provide utility improvements in database-dependent suppression, may incur, on the other, costs for the privacy guarantees. In particular, $\mathcal{S}$ needs to respect the neighborhood relation to ensure low privacy parameters [3]. We illustrate this phenomenon in our privacy results in Section 5.

In the following, we provide different answers to the suppression problem covering how suppression affects the privacy (Section 5) and utility guarantees (Section 6), ultimately seeing that our suppression does not improve over sampling.

5 The Effect of Suppression on Privacy

In this section, we study how the privacy guarantees are affected by suppression algorithms. As previously mentioned, we will assume that $\mathcal{M}$ satisfies unbounded (ϵ, δ) -DP, and study when $\mathcal{M} \circ \mathcal{S}$ satisfies unbounded $(\epsilon^{\mathcal{S}}, \delta^{\mathcal{S}})$ -DP, deriving expressions for $\epsilon^{\mathcal{S}}$ and $\delta^{\mathcal{S}}$. We will provide bounds independent of the choice of mechanism $\mathcal{M}$, and thus show the worst-case bounds with respect to $\mathcal{M}$.

Our goal is to show not only how specific targeted suppression, such as deleting outliers, affects the privacy parameters, but also the effect of any possible alternative suppression algorithm. This allows to identify which and how records can be deleted to improve the privacy guarantees, and thus we pave the way for data curators to easily learn what the privacy guarantees are after deleting exactly the records they want. To stay true to this, we want to impose the fewest conditions on $\mathcal{S}$ to provide the most general results possible.

We note that $\mathcal{S}$ does not necessarily satisfy (ϵ, δ) -DP and therefore the DP composition rules cannot be applied. Partic-

ularly for pure DP, since $\mathcal{S}$ only outputs subsets of the input database, it is possible that $\mathcal{S}(D)$ outputs a subset which cannot be output by $\mathcal{S}(D')$, violating the ϵ -DP definition. Therefore, the only algorithm $\mathcal{S}$ that satisfies pure DP is the mechanism that deletes all records, i.e., $\mathcal{S}(D) = \emptyset$ for all $D \in \mathbb{D}$, which satisfies 0-DP. We note that $\mathcal{M} \circ \mathcal{S}$ can still satisfy DP even when $\mathcal{S}$ is non-DP.

In the following, we first tackle deterministic suppression (Section 5.1), presenting some interesting results that show an initial understanding how deleting records affects the privacy parameters. Due to the limitations of this type of suppression in DP, we then study probabilistic suppression in a general way (Section 5.2) and provide a specific probabilistic suppression strategy for outlier deletion (Section 5.3).

We include proof sketches of all our results in Appendix A and the full detailed proofs in [arXiv:2601.05180](https://arxiv.org/abs/2601.05180).

5.1 Deterministic Suppression

It is well known that (ϵ, δ) -DP is a worst-case metric since any (ϵ, δ) -DP mechanism $\mathcal{M}$ must satisfy $\mathbb{P}\{\mathcal{M}(D) \in A\} \leq e^\epsilon \mathbb{P}\{\mathcal{M}(D') \in A\} + \delta$ for all neighboring $D, D' \in \mathbb{D} := \text{Domain}(\mathcal{M})$ (and all measurable $A \subseteq \text{Range}(\mathcal{M})$). One of the first results where we would intuitively expect privacy amplification is in reducing the number of inequalities that must be satisfied, which can easily be done by reducing the domain $\mathbb{D}$ of $\mathcal{M}$. In particular, by excluding the hardest-to-satisfy inequalities, we can lower the values of ϵ and δ , potentially obtaining that $\mathcal{M}$ over this reduced domain is (ϵ', δ') -DP with $\epsilon' < \epsilon$ and $\delta' < \delta$. We can obtain a domain reduction if we have a deterministic suppression algorithm that verifies $\mathcal{S}(\mathbb{D}) \subsetneq \mathbb{D}$, thus reducing the input of mechanism $\mathcal{M}$ from $\mathbb{D}$ to $\mathcal{S}(\mathbb{D})$.

Theorem 5.1 shows how the privacy parameters of $\mathcal{M}$ are affected when preprocessed with a deterministic suppression algorithm $\mathcal{S}$. There are two factors that affect the privacy parameters of $\mathcal{M} \circ \mathcal{S}$: First, the privacy improvement we can gain by restricting the domain of $\mathcal{M}$ as we explained; and second, the effect on the privacy parameters caused by applying $\mathcal{S}$, which closely follows from the known preprocessing result on c -stable transformations [31].

Theorem 5.1 (Effect of deterministic suppression). *Let $\mathcal{M}$ be an (ϵ, δ) -DP mechanism and $\mathcal{S}$ be a deterministic suppression algorithm, both with domain $\mathbb{D}$. Let $\mathbb{S}$ be such that $\mathcal{S}(\mathbb{D}) \subseteq \mathbb{S} \subseteq \mathbb{D}$, and suppose the restriction of $\mathcal{M}$ to domain $\mathbb{S}$, $\mathcal{M}|_{\mathbb{S}}$, is $(\epsilon_{\mathbb{S}}, \delta_{\mathbb{S}})$ -DP. Then, $\mathcal{M} \circ \mathcal{S} = \mathcal{M}|_{\mathbb{S}} \circ \mathcal{S}$ with domain $\mathbb{D}$ is $(\epsilon_{\mathbb{S}} \Delta_{\mathbb{S}} \mathcal{S}, \delta_{\mathbb{S}} \sum_{k=0}^{\Delta_{\mathbb{S}} \mathcal{S}-1} e^{\epsilon_{\mathbb{S}} k})$ -DP, where the sensitivity of $\mathcal{S}$ is*

$$\Delta_{\mathbb{S}} \mathcal{S} := \sup_{\substack{D, D' \in \mathbb{D} \\ \text{neighb.}}} d_{\mathbb{S}}(\mathcal{S}(D), \mathcal{S}(D')),$$

and $d_{\mathbb{S}}(\mathcal{S}(D), \mathcal{S}(D'))$ is the minimum number of neighboring databases in $\mathbb{S}$ needed to go from $\mathcal{S}(D)$ to $\mathcal{S}(D')$ (see [21]).

In the theorem, we show the effect of restricting the domain of $\mathcal{M}$ to intermediate subsets $\mathbb{S}$ (such that $\mathcal{S}(\mathbb{D}) \subseteq \mathbb{S} \subseteq \mathbb{D}$),

since it is possible that the subset $\mathbb{S}$ that provides the lowest privacy parameters is not $\mathcal{S}(\mathbb{D})$. Since the result holds for every choice of $\mathbb{S}$, we can choose $\mathbb{S}$ that minimizes the privacy parameters of $\mathcal{M} \circ \mathcal{S}$. We note that finding the minimum can be difficult since smaller $\mathbb{S}$ intuitively yield smaller (or equal) values of $\epsilon_{\mathbb{S}}$ and $\delta_{\mathbb{S}}$ but larger (or equal) values of $\Delta_{\mathbb{S}} \mathcal{S}$. More formally, given $\mathbb{S} \subseteq \mathbb{S}'$, we have that $\epsilon_{\mathbb{S}} \leq \epsilon_{\mathbb{S}'}$ and $\delta_{\mathbb{S}} \leq \delta_{\mathbb{S}'}$ (if chosen tightly), but, at the same time, we have $\Delta_{\mathbb{S}'} \mathcal{S} \leq \Delta_{\mathbb{S}} \mathcal{S}$. In particular, we note that the smallest sensitivity is

$$\Delta_{\mathbb{D}} \mathcal{S} = \sup_{\substack{D, D' \in \mathbb{D} \\ \text{neighb.}}} |\mathcal{S}(D) \Delta \mathcal{S}(D')| \leq \Delta_{\mathbb{S}} \mathcal{S}.$$

Moreover, there can exist $\mathbb{S} \subsetneq \mathbb{D}$ such that $\Delta_{\mathbb{D}} \mathcal{S} = \Delta_{\mathbb{S}} \mathcal{S}$ (e.g., $\mathbb{S} = \bigcup_{D \in \mathbb{D}} \mathcal{P}(\mathcal{S}(D))$ where $\mathcal{P}$ denotes the power set). In particular, we also note that if $\delta = 0$, then $\mathcal{M} \circ \mathcal{S}$ is $(\epsilon \Delta_{\mathbb{D}} \mathcal{S})$ -DP, remaining in pure DP.

We now provide an applied example of Theorem 5.1 and some of its consequences.

Example 5.2 (Laplace mechanism with deterministic suppression). Recall that the Laplace mechanism $\mathcal{M}_{f,b}$ of the query function $f: \mathbb{D} \rightarrow \mathbb{R}^k$ that adds noise drawn from the Laplace distribution $\text{Lap}(b)$ with scale b to each coordinate of $f(D)$ satisfies ϵ -DP with $\epsilon = \frac{\Delta f}{b}$ [17], where the sensitivity of f ,

$$\Delta f = \sup_{\substack{D, D' \in \mathbb{D} \\ \text{neighb.}}} \|f(D) - f(D')\|_1,$$

depends on $\mathbb{D}$ and thus on the range of f . Considering a deterministic suppression $\mathcal{S}$ that reduces the domain (i.e., $\mathcal{S}(\mathbb{D}) \subsetneq \mathbb{D}$) and selecting $\mathbb{S} := \bigcup_{C \in \mathbb{D}} \mathcal{P}(\mathcal{S}(C))$, we obtain by Theorem 5.1 that $\mathcal{M}_{f,b} \circ \mathcal{S}$ is $(\epsilon_{\mathbb{S}} \Delta_{\mathbb{S}} \mathcal{S})$ -DP with

$$\epsilon_{\mathbb{S}} = \frac{\Delta f|_{\mathbb{S}}}{b} \quad \text{and} \quad \Delta f|_{\mathbb{S}} := \sup_{\substack{D, D' \in \mathbb{S} \\ \text{neighb.}}} \|f(D) - f(D')\|_1 \leq \Delta f.$$

That is, we obtain $\epsilon_{\mathbb{S}} < \epsilon$ if and only if $\Delta f|_{\mathbb{S}} \Delta_{\mathbb{S}} \mathcal{S} < \Delta f$. This improvement can be leveraged to increase utility by raising the privacy parameter, which adds less noise (i.e., noise drawn from $\text{Lap}(b')$ with $b' < b$) accordingly. Selecting b' such that $\frac{\Delta f|_{\mathbb{S}}}{b'} \Delta_{\mathbb{S}} \mathcal{S} = \frac{\Delta f}{b}$ holds, will ensure that the privacy parameters of both mechanisms remains constant and allow us to evaluate the effect of suppression (cf. Section 6).

Furthermore, bear in mind that $\epsilon_{\mathbb{S}}$ and $\delta_{\mathbb{S}}$ depend on $\mathcal{M}$ and may change for different mechanisms. In particular, there are mechanisms that cannot benefit from a domain reduction, like a Laplace or Gaussian mechanism for a counting query (since $\Delta f|_{\mathbb{S}} = \Delta f = 1$ for all $\mathbb{S}$). Therefore, $\mathcal{M} \circ \mathcal{S}$ is always $(\epsilon \Delta_{\mathbb{S}} \mathcal{S}, \delta \sum_{k=0}^{\Delta_{\mathbb{S}} \mathcal{S}-1} e^{\epsilon k})$ -DP for all (ϵ, δ) -DP mechanisms $\mathcal{M}$, which is an independent bound on the choice of $\mathcal{M}$.

Moreover, Theorem 5.1 provides a tight bound: For all privacy parameters and all $\mathcal{S}$, there exists a DP mechanism

$\mathcal{M}$ such that $\mathcal{M} \circ \mathcal{S}$ is tightly $(\epsilon_{\mathbb{S}} \Delta_{\mathbb{S}} \mathcal{S}, \delta_{\mathbb{S}} \sum_{k=0}^{\Delta_{\mathbb{S}} \mathcal{S}-1} e^{\epsilon_{\mathbb{S}} k})$ -DP if $\delta_{\mathbb{S}} \sum_{k=0}^{\Delta_{\mathbb{S}} \mathcal{S}-1} e^{\epsilon_{\mathbb{S}} k} < 1$ (see Proposition A.1). Therefore, Theorem 5.1 provides a complete characterization for deterministic suppression, indicating that there are suppression algorithms $\mathcal{S}$ and mechanisms $\mathcal{M}$ such that $\mathcal{M} \circ \mathcal{S}$ provides weaker privacy than $\mathcal{M}$, since $\Delta_{\mathbb{S}} \mathcal{S}$ can potentially be greater than 1.

When we have a suppression algorithm with sensitivity $\Delta_{\mathbb{S}} \mathcal{S} = 1$, we obtain that the privacy parameters given by Theorem 5.1 remain constant (or decrease, if so by a domain reduction). By definition, all database-independent suppression algorithms $\mathcal{S}$ have sensitivity $\Delta_{\mathbb{S}} \mathcal{S} = 1$, such as fixing a subset A of the universe of records $\mathcal{X}$ and defining $\mathcal{S}_A(D) = D \cap A$ for all $D \in \mathbb{D}$. This can be understood as removing the values outside A , the set of elements with “good” properties, or the records that are not outlying. For instance, we can use this to remove predefined extreme values, such as super-centenarians in an age database, or remote locations in a location database. In both of these examples, the suppressed records are defined independently of the choice of database, i.e., using public or common knowledge to designate people over a certain age as outliers or to define which map areas are remote.

Alternatively, database-dependent suppression can be useful for outlier deletion because it does not require any knowledge and allows suppression on a per-database basis. As a simple example, consider the database class $\mathbb{D}$ of databases containing people’s ages (ranging from 0 to the maximum verified age) and other data. Applying a database-independent suppression that deletes all supercentenarians may make sense for many $D \in \mathbb{D}$, but the results can become skewed for specific databases in $\mathbb{D}$ such as a superagers database.

However, as covered in Section 4, DP must account for changes in-between databases, which increases the privacy parameters when applying a database-dependent suppression. In this case, this privacy degradation is represented by the sensitivity of $\mathcal{S}$, which is large or even infinite when defining a suppression strategy specifically to delete outliers or distant records. For example, deleting all records whose average distance to the other records in the database exceeds a certain threshold (Proposition A.2(a)) or deleting the top $P\%$ of records that are furthest away from all other records in the database (for $P \leq 50$; Proposition A.2(b)) are both suppression algorithms with $\Delta_{\mathbb{D}} \mathcal{S} = \infty$ (and thus $\Delta_{\mathbb{S}} \mathcal{S} = \infty \geq \Delta_{\mathbb{D}} \mathcal{S}$ for all $\mathbb{S}$). In these examples, adding or removing a record can have a large effect on the distance to the rest of the records in the database, and thus $\mathcal{S}(D)$ and $\mathcal{S}(D')$ can potentially be very different, which leads to $\Delta_{\mathbb{D}} \mathcal{S} = \infty$ and no DP guarantees.

In general, we find that many database-dependent suppression algorithms defined to suppress outliers require large or even infinite sensitivities, thus increasing the privacy parameters to unmanageable levels. Furthermore, while it is theoretically possible to construct a deterministic database-dependent suppression algorithm $\mathcal{S}$ with $\Delta_{\mathbb{D}} \mathcal{S} = 1$ (see Remark A.3), we have not found any that correspond to a meaningful way of deleting outliers.

5.2 Probabilistic Suppression

In this section, we consider probabilistic suppression and extend the privacy evaluation of suppression to the whole spectrum, covering any suppression algorithm. Precisely, we provide results showing how the privacy parameters of $\mathcal{M}$ are affected when preprocessing with any probabilistic suppression and when suppression yields privacy amplification.

Deterministic suppression can be viewed as a special case of probabilistic suppression, and thus our theorems generalize the results of the previous section. In particular, probabilistic suppression can still provide cases where $\mathcal{M} \circ \mathcal{S}$ is not DP, as we obtained in Section 5.1. Therefore, we find it convenient to exclude such cases from our main theorem (Theorem 5.3) in order to provide a concise result, and we later explain how we can extend the theorem to the rest of the suppression algorithms, including those that do not achieve DP.

Our proofs follow the steps of the existing theorem on Poisson sampling by Li et al. [29] but with our generalized suppression algorithm $\mathcal{S}$. We find that the newer sampling theorems [3, 39] require additional conditions that do not directly generalize, or are inapplicable, to the more general suppression. Like these sampling results, our theorems work for any (ϵ, δ) -DP mechanism $\mathcal{M}$ and the bounds are independent of the choice of $\mathcal{M}$. The bounds given are therefore worst-case with respect to $\mathcal{M}$, i.e., they represent bounds to the largest possible privacy parameters over all DP mechanisms.

The essential steps in the proof are bounding the privacy parameters of $\mathcal{M} \circ \mathcal{S}$ with that of $\mathcal{M}$ using the law of total probability, i.e.,

$$\mathbb{P}\{\mathcal{M}(\mathcal{S}(D)) \in A\} = \sum_{C \in \text{supp}(\mathcal{S}(D))} \mathbb{P}\{\mathcal{M}(C) \in A\} \mathbb{P}\{\mathcal{S}(D) = C\},$$

and finding a relation between the probability measures of $\mathcal{S}(D)$ and $\mathcal{S}(D')$ (for all neighboring $D, D' \in \mathbb{D}$). Here, finding a good relation is the challenging part. For the state-of-the-art uniform sampling [3, 29, 39], this relation is simply constant over the support of $\mathcal{S}(D)$, but the relation for general $\mathcal{S}$ can be hard to define and very complex in some exceptional cases.

Therefore, as mentioned, we find it convenient to exclude these hard cases for now by assuming that $\mathcal{S}$ satisfies the *support condition*: For all unbounded-neighboring $D, D' \in \mathbb{D}$ of the form $D' = D_{+y} := D \uplus \{y\}$ and all $C \subseteq D$, there is a non-zero chance that $\mathcal{S}(D)$ outputs C if and only if there is a non-zero chance that $\mathcal{S}(D')$ outputs C or C_{+y} ; or, formally, $C \in \text{supp}(\mathcal{S}(D))$ if and only if $C \in \text{supp}(\mathcal{S}(D'))$ or $C_{+y} \in \text{supp}(\mathcal{S}(D'))$. We can think of this condition as basically ensuring that if an output C is possible for $\mathcal{S}(D)$, then C or C_{+y} is also possible for $\mathcal{S}(D')$, which avoids dividing by 0 in Theorem 5.3.

Theorem 5.3 (Suppression theorem). *Let $\mathcal{M}$ be a mechanism that satisfies unbounded (ϵ, δ) -DP and $\mathcal{S}$ be a suppression algorithm that satisfies the support condition, both with domain*

$\mathbb{D}$. Then, $\mathcal{M} \circ \mathcal{S}$ is unbounded (ϵ^S, δ^S) -DP with

$$\epsilon^S = \sup_{\substack{D, D' \in \mathbb{D} \\ \text{neighb.}}} \epsilon_{D, D'}^S \quad \text{and} \quad \delta^S = \sup_{\substack{D, D' \in \mathbb{D} \\ \text{neighb.}}} \delta_{D, D'}^S,$$

where $\epsilon_{D, D'}^S$ and $\delta_{D, D'}^S$ are as follows: If $D' = D_{+y}$, then

$$\epsilon_{D, D'}^S = \max_{C \in \text{supp}(\mathcal{S}(D))} \frac{\mathbb{P}\{\mathcal{S}(D) = C\}}{\mathbb{P}\{\mathcal{S}(D') = C\} + e^{-\epsilon} \mathbb{P}\{\mathcal{S}(D') = C_{+y}\}}$$

and

$$\delta_{D, D'}^S = \delta \sum_{C \in \text{supp}(\mathcal{S}(D))} \frac{\mathbb{P}\{\mathcal{S}(D) = C\} e^{-\epsilon} \mathbb{P}\{\mathcal{S}(D') = C_{+y}\}}{\mathbb{P}\{\mathcal{S}(D') = C\} + e^{-\epsilon} \mathbb{P}\{\mathcal{S}(D') = C_{+y}\}},$$

and, since the values are not symmetric with respect to D, D' ,

$$\epsilon_{D', D}^S = \max_{C \in \text{supp}(\mathcal{S}(D))} \frac{\mathbb{P}\{\mathcal{S}(D') = C\} + e^{\epsilon} \mathbb{P}\{\mathcal{S}(D') = C_{+y}\}}{\mathbb{P}\{\mathcal{S}(D) = C\}}$$

and

$$\delta_{D', D}^S = \delta \sum_{C \in \text{supp}(\mathcal{S}(D))} \mathbb{P}\{\mathcal{S}(D') = C_{+y}\} = \delta \mathbb{P}\{y \in \mathcal{S}(D')\}.$$

Our theorem always satisfies $\delta^S \leq \delta$, but ϵ^S can be larger, equal, or smaller than ϵ . Therefore, our theorem does not always show a privacy amplification, but rather shows how the parameters change through a suppression algorithm, and gives us an intuition on how deleting records affects the privacy parameters. The bound we provide is not tight in general, but it is for some specific algorithms $\mathcal{S}$, such as Poisson sampling.

Even though Theorem 5.3 is not generally tight, it provides an intuition on how the privacy budget is affected by suppression. When the values of $\frac{\mathbb{P}\{\mathcal{S}(D)=C\}}{\mathbb{P}\{\mathcal{S}(D_{+y})=C\}}$ and $\frac{\mathbb{P}\{\mathcal{S}(D)=C\}}{\mathbb{P}\{\mathcal{S}(D_{+y})=C_{+y}\}}$ (or their inverses) remain small, $\epsilon_{D, D'}^S$ (or $\epsilon_{D', D}^S$) remains small; but in other cases, $\max\{\epsilon_{D, D'}^S, \epsilon_{D', D}^S\}$ takes on larger values, increasing the privacy parameters of $\mathcal{M} \circ \mathcal{S}$. In particular, we see at play the fact that DP must always factor in changes between neighboring databases, as mentioned in Section 4. For example, if we define $\mathcal{S}$ as probabilistically deleting those points that are furthest away from the mean, we must take into account the differences in distribution caused by adding any single record in any database.

Our theorem provides an upper bound on the privacy guarantees of every suppression algorithm satisfying the support condition. The complexity of the equations in Theorem 5.3 are just a consequence of the potential complexity of suppression algorithms. Therefore, Theorem 5.3 is more useful in evaluating specific suppression strategies, as we will do in Section 5.3. In addition, it also provides the tight bound for Poisson sampling [29] and our tight bound for deterministic suppression with $\Delta_{\mathbb{D}}\mathcal{S} = 1$ we provided in Theorem 5.1.

As mentioned earlier, we can adapt the proof of Theorem 5.3 to obtain the result for all suppression/sampling algorithms $\mathcal{S}$. The full result is given in Theorem A.5: Its idea

is to assign each subset C that does not satisfy the support condition with another term C^* that does, so that the bound can be defined. However, there are multiple ways to assign C^* to C , each giving a different bound of the privacy parameters. In this case, the privacy parameters increase for assignments with large $|CAC^*|$, and we lose the guarantee that $\delta^S \leq \delta$ or that $\mathcal{M} \circ \mathcal{S}$ is pure DP if $\mathcal{M}$ is pure DP from Theorem 5.3.

In summary, Theorem 5.3 (and the general Theorem A.5) tells us that some probabilistic suppression can provide privacy amplifications just as sampling does, even without the improvement provided by domain reduction; while other suppression strategies can end up with larger privacy parameters, especially if $\mathcal{S}$ varies significantly between neighboring databases. However, even though there are specific examples that increase the values of the privacy parameters (e.g., the deterministic ones), we are unable to provide proof of the tightness of the results. Finally, we note that these theorems are defined independently of the choice of $\mathcal{M}$, showing how privacy degrades in the worst case for any (ϵ, δ) -DP mechanism. Nevertheless, certain mechanisms $\mathcal{M}$ could provide better bounds, such as the improvement provided by domain reduction.

In summary, selective suppression strategies can easily violate acceptable privacy bounds, resulting in privacy degradation instead of amplification; while suppression methods that delete records more uniformly across databases guarantee lower (better) privacy parameters.

5.3 Distance-Based Probabilistic Suppression

In this section, we present a type of suppression strategy $\mathcal{S}$ to deal with the presence of outliers in databases. The privacy parameters of $\mathcal{M} \circ \mathcal{S}$ are obtained through Theorem 5.3.

In this suppression strategy, every record is suppressed independently—as in Poisson sampling—but with a probability proportional to how different they are from the other records in the database. By carefully measuring these differences, we obtain that any DP mechanism $\mathcal{M}$ preprocessed by our $\mathcal{S}$ is also DP and can derive the precise expression of its privacy parameters (Theorem 5.4). In addition, our probabilistic result avoids the large sensitivities of deterministic database-dependent suppression that we saw in Section 5.1, recalling that the deterministic version of this result is not DP ($\Delta_{\mathbb{D}}\mathcal{S} = \infty$, as seen in Proposition A.2(a)).

Formally, the difference between records is given through a (normalized) bounded distance $d: \mathcal{X} \times \mathcal{X} \rightarrow [0, 1]$ defined over the data universe $\mathcal{X}$ (from where the databases are drawn). We choose parameters $m, M \in (0, 1)$ with $m \leq M$ to control the extent to which the property of being an outlier is considered. Precisely, we take the (m, M) -transformation T of d , defined as $T(x, y) = (m + (M - m)d(x, y)) \in [m, M]$ for all $x, y \in \mathcal{X}$. Then, for any non-empty database $D \in \mathbb{D}$, we define the outlier-score function $\text{out}_D: D \rightarrow [m, M]$ over D (with respect to T) such that $\text{out}_D(x) = \frac{1}{|D|} \sum_{y \in D} T(x, y)$ for all $x \in D$.

Our suppression algorithm will delete every record x in D

independently with probability $\text{out}_D(x)$, the average distance to all elements in D . By definition, $\text{out}_D(x)$ is guaranteed to be between m and M , thus providing a lower and upper bound on the probability of a record being deleted. A large difference $M - m$ means that the suppression strategy discriminates strongly between inliers and outliers, while $M - m = 0$ provides the uniform Poisson sampling, where each record is deleted independently with the same probability m .

Theorem 5.4 shows the effect of this family of suppression algorithms, which we call *outlier-score suppression*, on the privacy parameters of $\mathcal{M} \circ \mathcal{S}$.

Theorem 5.4 (Outlier-score suppression). *Let $\mathcal{S}$ be the outlier-score suppression algorithm that independently deletes each record $x \in D$ with probability $\text{out}_D(x)$, i.e., $\mathcal{S}$ is defined so that*

$$\mathbb{P}\{\mathcal{S}(D) = C\} = \prod_{x \in C} (1 - \text{out}_D(x)) \prod_{x \in D \setminus C} \text{out}_D(x)$$

for all $D \in \mathbb{D}$ and all $C \subseteq D$. Then, if $\mathcal{M}$ is (ϵ, δ) -DP, we obtain that $\mathcal{M} \circ \mathcal{S}$ is (ϵ^S, δ^S) -DP where $\delta^S = \delta(1 - m)$ and

$$\epsilon^S = \max_{p \in [0,1]} \max\{l_1(p), l_2(p), l_3\}$$

up to an error² of $2 \cdot 10^{-7}$, where

$$l_1(p) = \ln(e^\epsilon - (e^\epsilon - 1)(pM + (1 - p)m)) \\ + p \frac{M}{m} + (1 - p) \frac{1 - m}{1 - (pM + (1 - p)m)} - 1$$

and

$$l_2(p) = \ln \left(e^\epsilon - (e^\epsilon - 1) \left(pM + (1 - p) \frac{(M + m) - pM}{2 - p} \right) \right) \\ + p \frac{M}{m} + (1 - p) \frac{1 - \frac{(M + m) - pM}{2 - p}}{1 - M} - 1$$

for all $p \in [0, 1]$; and $l_3 = -\ln(e^{-\epsilon} + (1 - e^{-\epsilon})M) + 1 - \frac{1 - M}{1 - m}$.

Note that ϵ^S depends only on ϵ and the constants m and M , and δ^S only on δ and m , and neither depend on the choice of mechanism $\mathcal{M}$ nor on the distance d . We plot the expression of ϵ^S with respect to m and M for some values of ϵ in Figure 2 and provide an interactive plot that allows computing the exact ϵ^S from the three constants³. We also provide a closed form of the precise values in Proposition A.7, and we note that the maximum is usually obtained when $p = 0$ or $p = 1$,

²The proof of Theorem 5.4 is computer assisted and verified up to an error of $2 \cdot 10^{-7}$ for every value of m and M in $\{0.01, 0.02, \dots, 0.98, 0.99\}$ (with $m \leq M$) and every value of ϵ in $\{0, 0.01, 0.02, \dots, 1.98, 1.99, 2\}$, $\{2.1, 2.2, \dots, 9.9, 10\}$, and $\{11, 12, \dots, 99, 100\}$. Computational power is used to check that the bound we provide matches the empirically optimized value (see Remark A.11 for more details). We conjecture the expression extends to all m and M , and to all $\epsilon \leq 100$ due to the continuity of ϵ^S .

³<https://www.desmos.com/calculator/hydpubdqtm>

which quickly simplifies the expression of ϵ^S . In these cases, our result is tight with respect to Theorem 5.3, but we cannot show whether it is tight in general (see Remark A.8 for further discussion on tightness). We note that the complex expression of ϵ^S in Theorem 5.4 is a consequence of our efforts to provide the tightest bounds on the result.

We observe, similar to Theorem 5.3, that we always have $\delta^S \leq \delta$, but ϵ^S can take values larger, equal, or smaller than ϵ . Lower values are obtained near the diagonal, with the blue regions in Figure 2 representing the values such that $\epsilon^S \leq \epsilon$, i.e., the values of m and M such that $\mathcal{S}$ provides a privacy amplification. In particular, Theorem 5.4 generalizes uniform Poisson sampling, which corresponds to the algorithms on the diagonal (i.e., $m = M$). As seen in the plots, ϵ^S increases as $m \rightarrow 0$ or $M \rightarrow 1$, with the limit values being ∞ . In addition, the ratio $\frac{\epsilon^S}{\epsilon}$ converges to 1 when $\epsilon \rightarrow \infty$.

Overall, outlier-score suppression can be used to suppress outliers with higher probability. We note that our definition of outlyingness is similar to previous definitions of record vulnerability used in DP [32]. However, our theorem shows that greatly differentiating outliers quickly increases the privacy parameters (see Figure 2).

5.4 The Impact of Suppression on Privacy

Suppression in DP can offer a privacy amplification, but not always: Since DP must protect any difference between neighboring databases, flexible suppression algorithms—such that $\mathcal{S}(D)$ and $\mathcal{S}(D')$ behave differently for neighboring databases $D, D' \in \mathbb{D}$ —greatly increase the privacy parameters.

Nevertheless, our theorems in this section not only represent this phenomenon, but also show precisely how suppression affects the privacy parameters. In particular, this effect is depicted in our theorems with the sensitivity $\Delta_{\mathcal{S}} \mathcal{S}$ in the deterministic case, and with the ratios $\frac{\mathbb{P}\{\mathcal{S}(D)=C\}}{\mathbb{P}\{\mathcal{S}(D+y)=C+y\}}$ and $\frac{\mathbb{P}\{\mathcal{S}(D)=C\}}{\mathbb{P}\{\mathcal{S}(D+y)=C+y\}}$ (and their inverses) in the probabilistic case. Furthermore, the difference between m and M in Theorem 5.4 also exhibits the same effect: The privacy parameters increase if outliers are more distinguishable (i.e., if $M - m$ increases).

In this sense, DP causes a conflict between ensuring manageable privacy parameters in $\mathcal{M} \circ \mathcal{S}$ and having flexibility between $\mathcal{S}$ across databases. For example, in the deterministic case, we have seen that database-independent suppression $\mathcal{S}$ ensures $\Delta_{\mathcal{S}} \mathcal{S} = 1$; hence, $\mathcal{M} \circ \mathcal{S}$ satisfies DP with at least the same privacy parameters than $\mathcal{M}$. On the other hand, the flexibility of database-deterministic suppression usually leads to very high sensitivities and privacy parameters. In some cases, these values are infinite (e.g., when records furthest from the database centroid are deleted), in which case $\mathcal{M} \circ \mathcal{S}$ cannot possibly be DP. Nevertheless, when $\Delta_{\mathcal{S}} \mathcal{S}$ is not too large, there can be a privacy amplification if the mechanism and the suppression algorithm allow for a domain reduction. However, this domain reduction is highly dependent on the

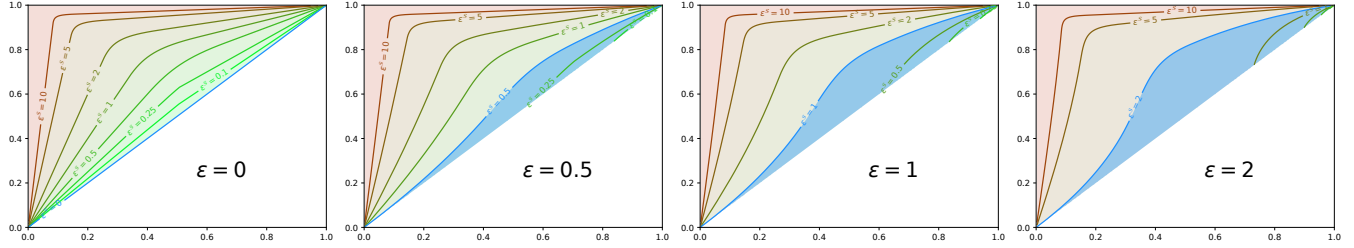


Figure 2: Contour plots of $\epsilon^S := \epsilon^S(\epsilon, m, M)$ with respect to m (x-axis) and M (y-axis) for the values of $\epsilon \in \{0, 0.5, 1, 2\}$. Values increase as $M - m$ increases. The blue regions show the values of m and M with a privacy amplification (i.e., such that $\epsilon^S \leq \epsilon$).

chosen mechanism, and there are always mechanisms for each suppression algorithm which do not benefit from it. Consequently, the guarantee here is not global.

In contrast to deterministic suppression, our results show that in the probabilistic case, privacy amplifications can exist independently of $\mathcal{M}$. However, they are only possible if the probability of deleting records does not vary drastically across neighboring databases; otherwise, privacy parameters degrade. We recall that suppression encompasses the state-of-the-art sampling, and it is especially the suppression algorithms close to uniform sampling that achieve a more significant privacy amplification—after all, these suppression algorithms ensure similar $\mathcal{S}(D)$ and $\mathcal{S}(D')$. In addition, the privacy parameters of suppression naturally cannot be less than those of uniform Poisson sampling (see Proposition A.9 and Corollary A.10).

In summary, due to the properties of DP, uniform suppression provides lower privacy parameters than targeted suppression. Our theorems confirm that privacy amplification is still achievable for other suppression strategies close to uniform suppression. However, the privacy parameters will easily start to increase when specifically protecting database outliers or vulnerable records if these records vary greatly between neighboring databases. This is especially true for deterministic suppression, which is more unforgiving than probabilistic suppression.

6 The Effect of Suppression on Utility

Having seen how suppression affects privacy, we are now interested in how it affects the mechanisms utility guarantees. In this section, we perform the same empirical evaluations we conducted for sampling in Section 3, that is, we compute the empirical evaluations of the utility guarantees of $\mathcal{M} \circ \mathcal{S}$ and compare them to those of $\mathcal{M}$.

In this case, we believe that we cannot formulate a fair evaluation using deterministic suppression: Any reasonable suppression strategy that is database-independent leads to a very large (or infinite) privacy budget for $\mathcal{M} \circ \mathcal{S}$, and database-dependent suppression requires defining global statistics about the database, which could lead to an unfair comparison since it requires knowledge about the values in the database.

Therefore, our evaluation centers on the outlier-score suppression algorithm of Section 5.3.

Note too that we only expect utility gains with a privacy amplification: If $\mathcal{S}$ actually increases the privacy parameters, we would generally obtain a noise amplification—not reduction—when translating the privacy parameters down to those of $\mathcal{M}$, and $\mathcal{M} \circ \mathcal{S}$ would expectedly provide worse utility than $\mathcal{M}$ at fixed privacy levels under our utility metrics.

6.1 Experiment Description and Setup

Conducting experiments similar to those in Section 3, we will compare the utility guarantees of $\mathcal{M} \circ \mathcal{S}$ to those of $\mathcal{M}$ under the same privacy guarantees. To ensure the same privacy level, we conduct the analogous transformation: By Theorem 5.4, if $\mathcal{M}$ satisfies (ϵ, δ) -DP, then $\mathcal{M} \circ \mathcal{S}$ satisfies (ϵ^S, δ^S) -DP with $\epsilon^S = \epsilon^S(\epsilon, m, M)$ and $\delta^S = \delta^S(\delta, m)$. So, to ensure that $\mathcal{M} \circ \mathcal{S}$ also satisfies (ϵ, δ) -DP, we impose $\mathcal{M}$ to be (ϵ'', δ'') -DP such that $\epsilon^S(\epsilon'', m, M) = \epsilon$ and $\delta^S(\delta'', m) = \delta$. We note that this process requires $\epsilon^S(\epsilon, m, M)$ to have an inverse ϵ'' with respect to ϵ , which is not possible when $\epsilon < \epsilon^S(0, m, M)$. This limitation is reflected by the unfilled areas in our plots.

We use the same mechanisms, utility metrics, databases, and privacy parameters ϵ and δ as in Section 3. As the mechanisms are randomized, we also compute $u(\mathcal{M}, D)$ and $u(\mathcal{M} \circ \mathcal{S}, D)$ the same amount of times as before, and provide their means. The only addition for this experiment is the distance function d for $\mathcal{S}$, directly linked to how records are suppressed. In this case, we select distances that intuitively represent ways of deleting records and mimic potential choices made by data curators.

Distance functions for $\mathcal{S}$. For the mean calculation, we select the absolute difference between values (i.e., the L^k distance in $\mathbb{R}$ for any $k \in \mathbb{N}$). Thus, the suppression algorithm deletes with higher probability the values that are the furthest away from others in a weighted manner. Note that the mean minimizes this average distance, so records closer to it are less likely to be deleted.

For the mode calculation, we select d as the discrete metric (i.e., $d(x, y) = 1$ if $x \neq y$ and $d(x, x) = 0$). This ensures that the values with higher counts in the database will be deleted with a lower probability than those with fewer.

For the clustering mechanisms, we choose d as the L^2 distance between records. This is the same distance function used in the mechanisms for the assignment of clusters/medians and in their respective utility metrics.

6.2 Experimental Results

To show whether $\mathcal{M} \circ \mathcal{S}$ can preserve utility better than $\mathcal{M}$, we plot the utility difference $u(\mathcal{M}, D) - u(\mathcal{M} \circ \mathcal{S}, D)$ (see Figure 3). Since higher values of $u(\mathcal{M}, D)$ are associated with worse utility, we obtain that $\mathcal{M} \circ \mathcal{S}$ provides better utility than $\mathcal{M}$ when $u(\mathcal{M}, D) - u(\mathcal{M} \circ \mathcal{S}, D)$, or the plot values, are positive. For each $\epsilon \in \{0.25, 0.5, 1, 2\}$, we provide plots of $u(\mathcal{M}, D) - u(\mathcal{M} \circ \mathcal{S}, D)$ with respect to m and M (similar to the privacy plots in Figure 2). We compute the difference value at the points where m and M both are in $\{0.1, \dots, 0.9\}$, with the precise value shown in the plots. The colors are then filled by triangulation, with the color grading being set to yellow for 0, red for negative values ($\mathcal{M} \circ \mathcal{S}$ provides worse utility), and green for positive values ($\mathcal{M} \circ \mathcal{S}$ provides better utility). The blue line corresponds to the values of m and M such that $\epsilon^S(\epsilon, m, M) = \epsilon$ (as in Figure 2).

All results across all databases and noise variations are qualitatively very similar, and we thus only plot representative examples in this section (Figure 3). All other plots are included as a gallery in the long version ([arXiv:2601.05180](https://arxiv.org/abs/2601.05180)).

Results. Our main observation is that there are almost no points where outlier-score suppression improves utility (cf., for instance, Figure 3). This holds even when the privacy amplification is compensated with lower noise (i.e., the values between the diagonal and the blue line in the plots).

In general, the figures also show how quickly privacy degrades as more and more data is suppressed, with some plot values largely increasing around $m = M = 0.9$. The difference also becomes smaller near $m = M = 0.1$, but we do not plot for smaller values because numerical inaccuracies begin to cause distortions. In the great majority of cases, the values of m and M that show the least utility loss are those on the diagonal, corresponding to uniform Poisson sampling (when comparing similar proportions of suppressed records). Our results also nearly always show that the utility difference under suppression worsens as ϵ increases (for fixed m and M).

We now look at the three main experiments individually. Figure 3 (top row) shows the results for the NoisyAverage mechanism with Laplace noise, which rarely sees a utility gain. However, the actual MPE difference remains at insignificantly low levels all throughout, around less than 0.6 p.p. (and sometimes even as low as 0.021 p.p.).

Similarly to the sampling counterpart, the utility of mode preservation decreases significantly when suppression is applied. As shown in Figure 3 (middle row), there is up to an 87.9 p.p. difference between the percentages of (in)correctly returning the mode of $\mathcal{M}$ and $\mathcal{M} \circ \mathcal{S}$ for RNM with Laplace noise on the age column (Adult database). These large differ-

ences are observed for values of m and M closer to 1, which correspond to deleting a large portion of the database. The utility values, however, strongly depend on the record distribution in the database, similar to our observations in Section 3.

The evaluation on clustering differs from the others, in the sense that the chosen utility metric does not assume the original database to be the ground truth or the base measurement. This means that the effect of suppression on the measured utility should theoretically be less damaging. Nevertheless, our results on k -median and DPLloyd show the same phenomenon as the mean computation: $\mathcal{M} \circ \mathcal{S}$ provides worse utility than $\mathcal{M}$, with the difference remaining small all throughout (see Figure 3, bottom row). DPLloyd shows some extremely small improvement over $\mathcal{M}$ for $m = M = 0.1$ for $\epsilon = 0.25$ and $\epsilon = 0.5$ values, but these could be due to rounding errors caused for these low parameters.

6.3 The Impact of Suppression on Utility

Finally, we investigated whether the privacy amplification achieved by outlier-score suppression could result in less perturbation of the DP mechanism $\mathcal{M}$ such that the overall utility of $\mathcal{M} \circ \mathcal{S}$ is greater than $\mathcal{M}$ at fixed privacy levels. Our results show that this is not generally possible: The utility loss from outlier-score suppression carries much more weight in the utility measurement than the DP perturbation, and almost all reductions in the DP perturbation achieved through privacy amplification are too low to benefit the tradeoff. In particular, this effect appears even in cases where the utility loss is insignificantly small, such as in the NoisyAverage experiments. Our results here thus largely follow the previous observations we made for sampling in Section 3. Moreover, we also note that most of our plots show a smaller utility difference on the diagonal. Hence, we see that utility is less affected by uniform Poisson sampling than by outlier-score suppression.

We believe that these results may be unexpected in some cases: Our method of assigning outliers depends on the chosen distances, which are deliberately selected to reduce privacy challenges and enhance utility. For instance, the distance chosen for the mode computation ensures that records with higher counts are less likely to be deleted than those with lower counts. However, in this case, the utility loss from deleting records is much greater than the utility gain from this selective deletion: Although the relative frequency of the mode in the database theoretically increases, the overall database size decreases, bringing the mode closer in count to the other records. This is why we obtain a significant loss of utility. Nevertheless, the fact that this type of suppression performs worse than uniform suppression, even when deleting the same proportion of records, offers an interesting insight into the privacy–utility tradeoff. We attribute this to the greater noise reduction in uniform suppression compared to non-uniform suppression due to the larger privacy amplification.

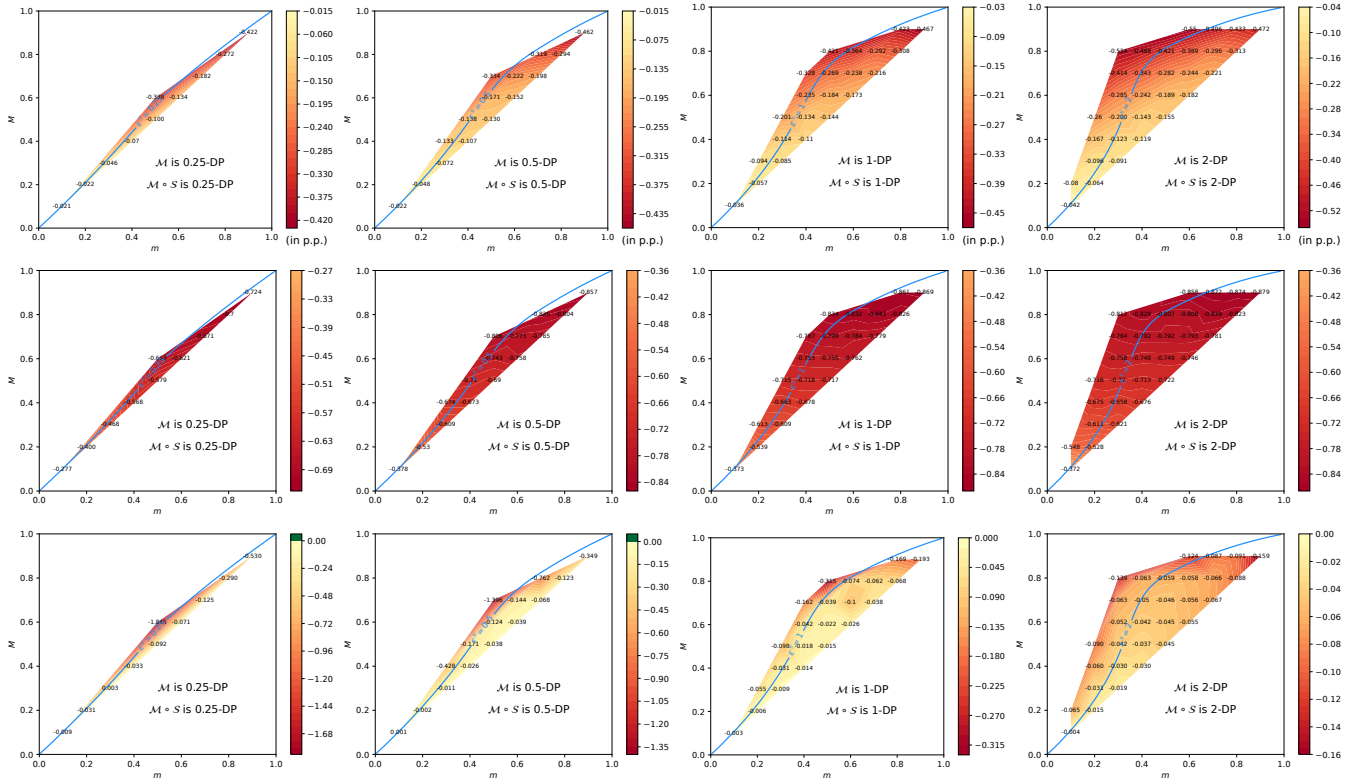


Figure 3: Plots of the utility difference of $\mathcal{M}$ minus that of $\mathcal{M} \circ \mathcal{S}$ (i.e., $u(\mathcal{M}, D) - u(\mathcal{M} \circ \mathcal{S}, D)$) for the Adult database (and age column). The shown mechanisms are (top row) NoisyAverage with Laplace mechanisms; (middle row) RNM with Laplace noise; and (bottom row) DPLloyd clustering algorithm.

7 Conclusion

Though sampling can provide orthogonal benefits like reducing time complexity, our privacy study shows that classic DP mechanisms without sampling consistently achieve a better privacy–utility tradeoff than the mechanisms with sampling, even when accounting for the noise reduction potentially gained by privacy amplification. These results motivated us to study the actual effect on the DP privacy–utility tradeoff when records are deleted as wanted. Here, we show that the positive suppression effects enjoyed in such privacy notions as k -anonymity do not transfer, mainly due to the DP definition itself. Since DP must account for any change between databases, more flexible strategic suppression algorithms that delete in a per-database setting come either at a weaker privacy amplification or, most often, at a privacy degradation.

We observe that our theorems in Section 5 describe the privacy amplification—or reduction—effect that any suppression strategy can enjoy. This is a new and rigorous insight into the effects that any conceivable way of records omission may exert on the privacy guarantees of DP mechanisms.

Our evaluation on outlier-score suppression yields the same results as in sampling. For all databases and mechanisms tested, we found that utility is reduced compared to analyses

without this preprocessing step at fixed privacy levels. However, the difference in error is quite insignificant in some cases. We conclude that, in our case, the potential utility gain from modifying the privacy parameters is insufficient to overcome the substantial utility loss caused by omitting records. In addition, since our tested mechanisms include the canonical building blocks of most DP mechanisms, we expect this result to extend to most of them. Furthermore, our results show that, in the majority of cases, uniform Poisson sampling provides the least utility loss among our tested suppression algorithms.

As future work, we will extend our results and evaluations to cover other suppression strategies and mechanisms and evaluate the effect of sampling and suppression over other DP variants, like bounded DP or Rényi DP. An interesting evaluation would be to see whether suppression can improve utility when measuring utility according to some ground truth rather than the input database.

In summary, our study provides new insights into sampling and suppression in DP, showing the particular need for balancing utility of these techniques against their demonstrated effects on the DP privacy–utility tradeoff. Overall, we show that sampling and outlier-score suppression both negatively impact the privacy–utility tradeoff, rendering the application of both techniques in DP questionable in this regard.

Ethical Considerations

The ethical implications of this work were thoroughly discussed. The authors of this paper declare the following:

Basic Ethical Principle and Stakeholder Analysis: This work was conducted following established ethical principles, guidelines, and best practices in data privacy. The focus is on differential privacy (DP), an important tool for protecting data with formal privacy guarantees. Our work deepens the understanding of sampling, a well-known technique in DP, and suppression, a technique used in privacy contexts but not in DP. Given the importance and popularity of DP, sampling, and suppression, our work may impact different members of society as it reveals vulnerabilities of these tools. In particular, we identify two main stakeholder groups: (1) individuals whose data were previously protected using a DP mechanism with sampling, and (2) data curators and practitioners who have developed or wish to develop mechanisms with sampling.

While our study shows that sampling and our suppression worsen utility, they do not compromise the level of privacy or protection they or any mechanism using them provides. Therefore, the individuals' privacy remains unaffected, regardless of our results. Our work does not cause any direct or indirect harm to the first stakeholder group or their data.

Our study may impact the second group, as it reveals a potential weakness in their mechanisms. While our work raises questions about the benefits of sampling, it does not weaken or invalidate published mechanisms that use sampling. Rather, it reveals the possibility of refining the privacy–utility tradeoff, which may encourage data practitioners to conduct new experiments on previously published work. Similarly, our study emphasizes the need to further evaluate the effect of sampling when designing new sampling-based mechanisms.

Bias and Fairness: We acknowledge the potential biases that can accompany any empirical evaluation. To mitigate these biases, we experimented with multiple mechanisms, databases, and parameters, and performed multiple iterations to guarantee statistical accuracy. We paid careful attention to avoid undue generalization of our results. Our experimentation and results only cover the privacy–utility tradeoff of sampling and our suppression in unbounded DP. We acknowledge that our conclusions may differ for other scenarios, suppression algorithms, or DP variants. We reiterate the need to evaluate the actual effect of sampling (or any preprocessing).

Decision: Our work raises awareness of sampling and suppression, improving their understanding of these techniques and discussing how they do not improve the privacy–utility tradeoff of DP mechanisms in our settings. Our paper also opens the door to testing mechanisms that have previously used sampling, potentially leading to improvements. Although our results are negative, we hope that they can guide the future design of more effective privacy-preserving mechanisms. The potential improvements that follow from our work have led us to submit it for consideration at USENIX.

Subjects and Public Database Use: Our work does not involve any direct human or animal participation. For our experiment, we used either synthetic databases or commonly used, publicly available standard databases derived from the US Census [6, 10] and Irish Census Data [2]. These databases had all been previously and thoroughly anonymized to prevent the identification of any individual and are explicitly released for public research, and statistical and educational use. In our work, we use these databases only for statistical research and do not intend to deanonymize them or perform data linking that might risk identifying the participants or their data.

Competing Interests: The authors have no relevant financial or nonfinancial interests to disclose, and are not affiliated with any organization or entity with a financial or nonfinancial interest in the content or implications of this paper.

Further ethical considerations: The research activities conducted in this paper did not have the potential to negatively affect the authors or their health. This work has been conducted within the framework of the law.

Open Science

To promote transparency and reproducibility, we make available an artifact that includes all the code from our work at: [10.5281/zenodo.17977527](https://doi.org/10.5281/zenodo.17977527). This includes the code for all the experiments (i.e., the mean and mode computation, and the two clustering mechanisms) for both sampling and suppression, as discussed in Sections 3 and 6, the code for the numerical computation for Theorem 5.4 (see Remark A.11), and the code to generate all plots in the paper. All databases used are included in the artifact because they are from publicly accessible data. We provide all the necessary code to reproduce every result and plot of our paper.

The long version of this paper, which includes the proofs of all our theorems and the plot gallery for all experiments, is available at [arXiv:2601.05180](https://arxiv.org/abs/2601.05180).

Acknowledgments

Javier Parra-Arnau is a “Ramón y Cajal” fellow (ref. RYC2021-034256-I) funded by the MCIN/AEI/10.13039/501100011033 and the EU “NextGenerationEU”/PRTR. This work was also supported by (i) the “DIstributed Smart Communications with Verifiable EneRgy-optimal Yields (DISCOVERY)” (PID2023-148716OB-C32) and the “Privacy-preserving techniques for mobility data” (CNS2025-166870) projects, funded by the same two institutions above; (ii) the Generalitat de Catalunya, under the AGAUR grant “2021 SGR 01413”; (iii) KASTEL Security Research Labs, Karlsruhe, funded by the Topic Engineering Secure Systems of the Helmholtz Association (HGF); and (iv) Germany’s Excellence Strategy – EXC 2050/1 – Project ID 390696704 – Cluster of Excellence “Centre for Tactile Internet with Human-

in-the-Loop” (CeTI). This paper has been edited by our textician, Daniel Shea. We would also like to thank Daniel Schadt, Patricia Guerra-Balboa, and Felix Morsbach for their useful help and comments.

References

- [1] Martín Abadi, Andy Chu, Ian Goodfellow, H. Brendan McMahan, Ilya Mironov, Kunal Talwar, and Li Zhang. Deep learning with differential privacy. In *Proc. ACM SIGSAC Conf. Comput., Commun. Secur. (CCS)*, 2016.
- [2] Vanessa Ayala-Rivera, A. Omar Portillo-Dominguez, Liam Murphy, and Christina Thorpe. COCOA: A synthetic data generator for testing anonymization techniques. In *Priv. Stat. Databases (PSD)*, 2016.
- [3] Borja Balle, Gilles Barthe, and Marco Gaboardi. Privacy amplification by subsampling: Tight analyses via couplings and divergences. In *Proc. Int. Conf. Neural Inform. Process. Syst. (NIPS)*, 2018.
- [4] Borja Balle and Yu-Xiang Wang. Improving the Gaussian mechanism for differential privacy: Analytical calibration and optimal denoising. In *Proc. Int. Conf. Mach. Learn. (ICML)*, 2018.
- [5] R.J. Bayardo and Rakesh Agrawal. Data privacy through optimal k-anonymization. In *Proc. IEEE Int. Conf. Data Eng. (ICDE)*, 2005.
- [6] Barry Becker and Ronny Kohavi. Adult, 1996.
- [7] Amos Beimel, Kobbi Nissim, and Uri Stemmer. Characterizing the sample complexity of private learners. In *Proc. Conf. Innov. Comput. Sci. (ITCS)*, 2013.
- [8] Jeremiah Blocki, Elena Grigorescu, and Tamalika Mukherjee. Differentially-private sublinear-time clustering. In *Proc. IEEE Int. Symp. Inform. Theory (ISIT)*, 2021.
- [9] Avrim Blum, Cynthia Dwork, Frank McSherry, and Kobbi Nissim. Practical privacy: The SuLQ framework. In *Proc. ACM Symp. Prin. Database Syst. (PODS)*, 2005.
- [10] Ruth Brand, Josep Domingo-Ferrer, and Josep M Mateo-Sanz. Reference data sets to test and compare SDC methods for protection of numerical microdata. Tech. Rep. European project IST-2000-25069 CASC, 2002.
- [11] Mark Bun, Jörg Drechsler, Marco Gaboardi, Audra McMillan, and Jayshree Sarathy. Controlling privacy loss in sampling schemes: An analysis of stratified and cluster sampling, 2023.
- [12] Kamalika Chaudhuri and Nina Mishra. When random sampling preserves privacy. In *Proc. Adv. Cryptology – Annual Int. Cryptology Conf. (CRYPTO)*, 2006.
- [13] Rui Chen, Qian Xiao, Yu Zhang, and Jianliang Xu. Differentially private high-dimensional data publication via sampling-based inference. In *Proc. ACM SIGKDD Int. Conf. Knowl. Disc., Data Min. (KDD)*, 2015.
- [14] Damien Desfontaines and Balázs Pejó. SoK: Differential privacies. *Proc. Priv. Enhanc. Technol.*, 2020.
- [15] Cynthia Dwork. Differential privacy. In *Proc. Int. Colloq. Automata, Lang., Program. (ICALP)*, 2006.
- [16] Cynthia Dwork, Krishnaram Kenthapadi, Frank McSherry, Ilya Mironov, and Moni Naor. Our data, ourselves: Privacy via distributed noise generation. In *Proc. Adv. Cryptology (EUROCRYPT)*, 2006.
- [17] Cynthia Dwork and Aaron Roth. The algorithmic foundations of differential privacy. *Found. Trends Theor. Comput. Sci.*, 2014.
- [18] Khaled El Emam and Fida Kamal Dankar. Protecting privacy using k-anonymity. *J. Amer. Med. Inform. Assoc.*, 2008.
- [19] Juanru Fang and Ke Yi. Privacy amplification by sampling under user-level differential privacy. *Proc. ACM Manag. Data*, 2024.
- [20] Marco Gramaglia, Marco Fiore, Angelo Furno, and Razvan Stanica. GLOVE: Towards privacy-preserving publishing of record-level-truthful mobile phone trajectories. *ACM/IMS Trans. Data Sci.*, 2021.
- [21] Patricia Guerra-Balboa, Àlex Miranda-Pascual, Javier Parra-Arnau, and Thorsten Strufe. Composition in differential privacy for general granularity notions. In *Proc. IEEE Comput. Security Found. Symp. (CSF)*, 2024.
- [22] Anupam Gupta, Katrina Ligett, Frank McSherry, Aaron Roth, and Kunal Talwar. Differentially private combinatorial optimization. In *Proc. ACM-SIAM Symp. Discr. Alg. (SODA)*, 2010.
- [23] Anco Hundepool, Josep Domingo-Ferrer, Luisa Francioni, Sarah Giessing, Eric S. Nordholt, Keith Spicer, and Peter-Paul de Wolf. *Statistical Disclosure Control*. Wiley, 2012.
- [24] Yangfan Jiang, Xinjian Luo, Yin Yang, and Xiaokui Xiao. Calibrating noise for group privacy in subsampled mechanisms, 2024.
- [25] Daniel Kifer and Ashwin Machanavajjhala. No free lunch in data privacy. In *Proc. ACM SIGMOD Int. Conf. Manage. Data (MOD)*, 2011.

- [26] Antti Koskela and Tejas D. Kulkarni. Practical differentially private hyperparameter tuning with subsampling. In *Proc. Int. Conf. Neural Inform. Process. Syst. (NeurIPS)*, 2023.
- [27] Jonathan Lebensold, Doina Precup, and Borja Balle. On the privacy of selection mechanisms with Gaussian noise. In *Proc. Int. Conf. Artif. Intell., Stat. (AISTATS)*, 2024.
- [28] Ninghui Li, Min Lyu, Dong Su, and Weining Yang. *Differential Privacy: From Theory to Practice*. Morgan & Claypool, 2016.
- [29] Ninghui Li, Wahbeh Qardaji, and Dong Su. On sampling, anonymization, and differential privacy or, k -anonymization meets differential privacy. In *Proc. ACM Int. Symp. Inform. Comput. Commun. Secur. (ASIACCS)*, 2012.
- [30] Bing-Rong Lin, Ye Wang, and Shantanu Rane. On the benefits of sampling in privacy preserving statistical analysis on distributed databases, 2013.
- [31] Frank McSherry. Privacy integrated queries. In *Proc. ACM SIGMOD Int. Conf. Manage. Data (MOD)*, 2009.
- [32] Matthieu Meeus, Florent Guepin, Ana-Maria Crețu, and Yves-Alexandre de Montjoye. Achilles’ heels: Vulnerable record identification in synthetic data publishing. In *Proc. Eur. Symp. Res. Comput. Secur. (ESORICS)*, 2024.
- [33] Ilya Mironov, Kunal Talwar, and Li Zhang. Rényi differential privacy of the sampled Gaussian mechanism, 2019.
- [34] Natalia Ponomareva, Hussein Hazimeh, Alex Kurakin, Zheng Xu, Carson Denison, H. Brendan McMahan, Sergei Vassilvitskii, Steve Chien, and Abhradeep Guha Thakurta. How to DP-fy ML: A practical guide to machine learning with differential privacy. *J. Artif. Intell. Res.*, 2023.
- [35] Ossi Räisä, Joonas Jälkö, and Antti Honkela. Subsampling is not magic: Why large batch sizes work for differentially private stochastic optimisation. In *Proc. Int. Conf. Mach. Learn. (ICML)*, 2024.
- [36] Pierangela Samarati. Protecting respondents identities in microdata release. *IEEE Trans. Knowl. Data Eng.*, 2001.
- [37] Adam Smith. Differential privacy and the secrecy of the sample, September 2009.
- [38] Jordi Soria-Comas, Josep Domingo-Ferrer, David Sánchez, and Sergio Martínez. Enhancing data utility in differential privacy via microaggregation-based k -anonymity. *VLDB J.*, 2014.
- [39] Thomas Steinke. Composition of differential privacy & privacy amplification by subsampling, 2022.
- [40] Dong Su, Jianneng Cao, Ninghui Li, Elisa Bertino, Min Lyu, and Hongxia Jin. Differentially private k -means clustering and a hybrid approach to private optimization. *ACM Trans. Priv. Secur.*, 2017.
- [41] Jonathan Ullman. CS7880: Rigorous approaches to data privacy, Spring 2017 POTW #1. 2017.
- [42] Yu-Xiang Wang, Jing Lei, and Stephen E. Fienberg. Learning with differential privacy: Stability, learnability and the sufficiency and necessity of ERM principle. *J. Mach. Learn. Res.*, 2016.
- [43] Yang Xu, Tinghuai Ma, Meili Tang, and Wei Tian. A survey of privacy preserving data publishing using generalization and suppression. *Appl. Math. Inf. Sci.*, 2014.
- [44] Yuqing Zhu and Yu-Xiang Wang. Poisson subsampled Rényi differential privacy. In *Proc. Int. Conf. Mach. Learn. (ICML)*, pages 7634–7642. PMLR, May 2019.

A Proof Sketches and Additional Results

A.1 Section 5.1: Deterministic Suppression

Proof sketch of Theorem 5.1. We fix $D, D' \in \mathbb{D}$. Suppose that $k = d_{\mathbb{S}}(\mathcal{S}(D), \mathcal{S}(D'))$ is finite (otherwise, $\Delta_{\mathbb{S}}\mathcal{S} = \infty$ and we are done). Then, there exists a chain of $k + 1$ databases $D_0, \dots, D_k \in \mathbb{S}$ such that consecutive databases are neighboring, and $D_0 = \mathcal{S}(D)$ and $D_k = \mathcal{S}(D')$. By repeatedly applying the definition of DP over $\mathcal{M}|_{\mathbb{S}}$, we obtain for all measurable $A \subseteq \text{Range}(\mathcal{M})$,

$$\begin{aligned} \mathbb{P}\{\mathcal{M}|_{\mathbb{S}}(\mathcal{S}(D)) \in A\} &\leq e^{\varepsilon_{\mathbb{S}}} \mathbb{P}\{\mathcal{M}|_{\mathbb{S}}(D_1) \in A\} + \delta_{\mathbb{S}} \leq \dots \\ &\leq e^{k\varepsilon} \mathbb{P}\{\mathcal{M}|_{\mathbb{S}}(\mathcal{S}(D')) \in A\} + \delta_{\mathbb{S}} \sum_{i=0}^{k-1} e^{\varepsilon_{\mathbb{S}} i}. \end{aligned} \quad (\text{A.1})$$

Finally, by applying $k \leq \Delta_{\mathbb{S}}\mathcal{S}$, we obtain the result. $\square$

Proposition A.1. *The bound provided by Theorem 5.1 is tight for all suppression algorithms $\mathcal{S}$ if $\delta_{\mathbb{S}} \sum_{i=0}^{\Delta_{\mathbb{S}}\mathcal{S}-1} e^{\varepsilon_{\mathbb{S}} i} < 1$.*

Proof sketch. The proof is a corollary of the existence of a mechanism $\mathcal{M}$ such that the group privacy bound is tight (i.e., all the inequalities in Equation (A.1) are always achieved by one mechanism $\mathcal{M}$). It suffices to select such a mechanism (e.g., the Laplace mechanism). $\square$

Proposition A.2. *Let $\mathbb{D} = \mathbb{D}_X$ be the class of all databases with elements drawn from X with $|X| \geq 2$. Let $d: X \times X \rightarrow [0, 1]$ be any distance over X and we denote $\text{avg}(x, D)$ the average distance of x to all elements in D .*

- (a) Consider the suppression strategy S_K such that $S_K(D) = \{x \in D \mid \text{avg}(x, D) \leq K\}$ for any $K \in (0, 1) \cap \mathbb{Q}$. Then S_K has sensitivity $\Delta_{\mathbb{D}} S_K = \infty$.
- (b) For any $p \in (0, \frac{1}{2}]$, consider the suppression strategy S_p that deletes the top $\lfloor p|D| \rfloor$ records with the highest average distance to the records of its input D (in the case of ties among the top $\lfloor p|D| \rfloor$ elements, we delete all elements). Then S_p has sensitivity $\Delta_{\mathbb{D}} S_p = \infty$.

Proof sketch. This proof consists of showing that there are neighboring databases, D and D' , such that many elements are deleted in $S(D)$ but not in $S(D')$. If this difference scales with the size of the databases, then the sensitivity must be infinite because it captures the worst-case difference. We will now describe the database pairs (D_n, D'_n) for which the difference becomes unbounded as $n \rightarrow \infty$.

(a) Since $K \in (0, 1) \cap \mathbb{Q}$, there exists $N \in \mathbb{N}, N \geq 2$ such that $NK \in \mathbb{N}$. We define the database D_n as having nNK copies of x'_n and $nN(1 - K)$ copies of y'_n (with $d(x'_n, y'_n) > 1 - \frac{1}{n}$), and we select D'_n as having one fewer x'_n . Then, we obtain both $\text{avg}(y'_n, D_n) \leq K$ and $\text{avg}(y'_n, D'_n) > K$. Consequently, S_K suppresses all copies of y'_n from D'_n but none from D_n . Therefore, for all $n \in \mathbb{N}$, these $nN(1 - K)$ elements are in $S_K(D_n) \Delta S_K(D'_n)$, and thus, $\Delta_{\mathbb{D}} S_K \geq \sup_{n \in \mathbb{N}} nN(1 - K) = \infty$.

(b) For all $N \in \mathbb{N}$ such that $\lfloor pN \rfloor \geq 1$, we consider the database D_N with N elements: $\lfloor pN \rfloor$ copies of x' and $N - \lfloor pN \rfloor$ copies of y' (with $x' \neq y'$). If $p \leq \frac{1}{2}$, then $\text{avg}(y', D_N) \leq \text{avg}(x', D_N)$. Therefore, the $\lfloor pN \rfloor$ records in D_N with the highest average distance are the $\lfloor pN \rfloor$ copies of x' . Thus, $S_p(D_N) = \{y', \dots, y'\}$, i.e., the $N - \lfloor pN \rfloor$ copies of y' .

Now, if we take D'_N with one fewer copy of x' , then we also obtain that $\text{avg}(y', D_N) \leq \text{avg}(x', D_N)$. However, since there are fewer than $\lfloor pN \rfloor$ copies of x' , every copy of y' ties as the $\lfloor pN \rfloor$ th element with the highest average distance. Consequently, by the definition of S_p , we obtain that $S_p(D'_N) = \emptyset$. Thus, $\Delta_{\mathbb{D}} S_p \geq \sup_{N \in \mathbb{N}, \lfloor pN \rfloor \geq 1} (N - \lfloor pN \rfloor) = \infty$. $\square$

Remark A.3. The suppression algorithm that, given D , outputs D if $|D| = 1$ and $\emptyset$ otherwise is data-dependent with sensitivity 1 by construction. This suppression algorithm has no applications and is only presented for illustrative purposes.

A.2 Section 5.2: Probabilistic Suppression

Proof sketch of Theorem 5.3. We fix $D, D' \in \mathbb{D}$ neighboring databases with $D' = D_{+y}$. To obtain the result, it is necessary that both of these inequalities hold for all measurable $A \subseteq \text{Range}(\mathcal{M} \circ S) \subseteq \text{Range}(\mathcal{M})$:

$$\mathbb{P}\{\mathcal{M}(S(D)) \in A\} \leq e^{\varepsilon_{D,D'}} \mathbb{P}\{\mathcal{M}(S(D')) \in A\} + \delta_{D,D'}^S,$$

$$\mathbb{P}\{\mathcal{M}(S(D')) \in A\} \leq e^{\varepsilon_{D',D}} \mathbb{P}\{\mathcal{M}(S(D)) \in A\} + \delta_{D',D}^S,$$

with $\varepsilon_{D,D'}^S, \delta_{D,D'}^S, \varepsilon_{D',D}^S$ and $\delta_{D',D}^S$ as defined in the statement. Both inequalities are proven using the law of total probability,

i.e.,

$$\mathbb{P}\{\mathcal{M}(S(D)) \in A\} = \sum_{C \in \text{supp}(S(D))} \mathbb{P}\{\mathcal{M}(C) \in A\} \mathbb{P}\{S(D) = C\},$$

and applying the DP property of $\mathcal{M}$ for the neighboring databases C and C_{+y} . Since the proofs of both inequalities follow the same principle, we include only the second one here for simplicity. Thus, by the law of total probability and further manipulations, we obtain that

$$\begin{aligned} \mathbb{P}\{\mathcal{M}(S(D')) \in A\} &= \sum_{C \in \text{supp}(S(D))} \mathbb{P}\{\mathcal{M}(C) \in A\} \mathbb{P}\{S(D') = C\} \\ &+ \sum_{C \in \text{supp}(S(D))} \mathbb{P}\{\mathcal{M}(C_{+y}) \in A\} \mathbb{P}\{S(D') = C_{+y}\}, \end{aligned}$$

where the equality requires the support hypothesis, i.e., $\text{supp}(S(D)) = \{C \subseteq D \mid C \text{ or } C_{+y} \text{ is in } \text{supp}(S(D'))\}$.

Now, applying the (ε, δ) -DP property to $\mathcal{M}$, we have that

$$\begin{aligned} \mathbb{P}\{\mathcal{M}(S(D')) \in A\} &\leq \sum_{C \in \text{supp}(S(D))} \mathbb{P}\{\mathcal{M}(C) \in A\} \mathbb{P}\{S(D') = C\} \\ &+ \sum_{C \in \text{supp}(S(D))} (e^\varepsilon \mathbb{P}\{\mathcal{M}(C) \in A\} + \delta) \mathbb{P}\{S(D') = C_{+y}\} \\ &= \underbrace{\sum_{C \in \text{supp}(S(D))} \mathbb{P}\{\mathcal{M}(C) \in A\} q_C}_{=:a} + \underbrace{\delta \sum_{C \in \text{supp}(S(D))} \mathbb{P}\{S(D') = C_{+y}\}}_{=:b} \end{aligned}$$

with $q_C := \mathbb{P}\{S(D') = C\} + e^\varepsilon \mathbb{P}\{S(D') = C_{+y}\}$. By further manipulating the values a and b , we obtain

$$\begin{aligned} a &= \sum_{C \in \text{supp}(S(D))} \mathbb{P}\{\mathcal{M}(C) \in A\} \mathbb{P}\{S(D) = C\} \frac{q_C}{\mathbb{P}\{S(D) = C\}} \\ &\leq \left(\max_{C \in \text{supp}(S(D))} \frac{q_C}{\mathbb{P}\{S(D) = C\}} \right) \sum_{C \in \text{supp}(S(D))} \mathbb{P}\{\mathcal{M}(S(D)) \in A\}, \end{aligned}$$

where the maximum is $e^{\varepsilon_{D',D}^S}$, and $b = \delta \mathbb{P}\{y \in S(D')\} = \delta_{D',D}^S$. We thus obtain the second inequality. The result follows by taking the supremum over neighboring D and D' . $\square$

Definition A.4 (Assignment function). Let S be a suppression algorithm with domain $\mathbb{D}$. Let $D, D' \in \mathbb{D}$ be unbounded-neighboring and assume $D' = D_{+y}$. Let $\text{ext}_D(S(D')) = \{C \subseteq D \mid C \in \text{supp}(S(D')) \text{ or } C_{+y} \in \text{supp}(S(D'))\}$.

We define an assignment $\mathcal{A}_{D,D'}$ from D to D' as any function $\mathcal{A}_{D,D'} : \text{supp}(S(D)) \rightarrow \text{ext}_D(S(D'))$, and an assignment $\mathcal{A}_{D',D}$ from D' to D as any function $\mathcal{A}_{D',D} : \text{ext}_{D'}(S(D')) \rightarrow \text{supp}(S(D))$.

We denote the fiber of C under $\mathcal{A}_{D,D'}$ by $\mathcal{A}_{D,D'}^{-1}[C]$, which is defined as the set of all databases C^* in the domain of $\mathcal{A}_{D,D'}$ such that $\mathcal{A}_{D,D'}(C^*) = C$ (analogously for $\mathcal{A}_{D',D}$). Naturally, if C is not in the image of $\mathcal{A}_{D,D'}$, we have that $\mathcal{A}_{D,D'}^{-1}[C] = \emptyset$.

Theorem A.5 (Suppression theorem (general result)). Let $\mathcal{M}$ with domain $\mathbb{D}$ be a mechanism that satisfies unbounded (ε, δ) -DP and S be a suppression algorithm with domain $\mathbb{D}$.

Then, $\mathcal{M} \circ \mathcal{S}$ is unbounded (ϵ^S, δ^S) -DP with

$$\epsilon^S = \sup_{\substack{D, D' \in \mathbb{D} \\ \text{neighb.}}} \epsilon_{D, D'}^S \quad \text{and} \quad \delta^S = \sup_{\substack{D, D' \in \mathbb{D} \\ \text{neighb.}}} \delta_{D, D'}^S,$$

where $\epsilon_{D, D'}^S$ and $\delta_{D, D'}^S$ depend on an assignment $\mathcal{A}_{D, D'}$ and are defined as follows: If $D' = D_{+y}$, then

$$\epsilon_{D, D'}^S = \max_{C \in \text{ext}_D(\mathcal{S}(D'))} \frac{\sum_{C^* \in \mathcal{A}_{D, D'}^{-1}[C]} e^{\epsilon |C \Delta C^*|} \mathbb{P}\{\mathcal{S}(D) = C^*\}}{\mathbb{P}\{\mathcal{S}(D') = C\} + e^{-\epsilon} \mathbb{P}\{\mathcal{S}(D') = C_{+y}\}}$$

and

$$\delta_{D, D'}^S = \delta \sum_{C \in \text{ext}_D(\mathcal{S}(D'))} \sum_{C^* \in \mathcal{A}_{D, D'}^{-1}[C]} \mathbb{P}\{\mathcal{S}(D) = C^*\} \cdot \left(\frac{e^{\epsilon(|C \Delta C^*| - 1)} \mathbb{P}\{\mathcal{S}(D') = C_{+y}\}}{\mathbb{P}\{\mathcal{S}(D') = C\} + e^{-\epsilon} \mathbb{P}\{\mathcal{S}(D') = C_{+y}\}} + \sum_{k=0}^{|C \Delta C^*| - 1} e^{\epsilon k} \right),$$

and, since the values are not symmetric with respect to D and D' ,

$$\epsilon_{D', D}^S = \max_{C \in \text{supp}(\mathcal{S}(D))} \left(\frac{1}{\mathbb{P}\{\mathcal{S}(D) = C\}} \sum_{C^* \in \mathcal{A}_{D', D}^{-1}[C]} e^{\epsilon |C \Delta C^*|} \cdot (\mathbb{P}\{\mathcal{S}(D') = C^*\} + e^{\epsilon} \mathbb{P}\{\mathcal{S}(D') = C_{+y}^*\}) \right)$$

and

$$\delta_{D', D}^S = \delta \sum_{C \in \text{supp}(\mathcal{S}(D))} \sum_{C^* \in \mathcal{A}_{D', D}^{-1}[C]} \left(\mathbb{P}\{\mathcal{S}(D') = C^*\} \sum_{k=0}^{|C \Delta C^*| - 1} e^{\epsilon k} + \mathbb{P}\{\mathcal{S}(D') = C_{+y}^*\} \sum_{k=0}^{|C \Delta C^*|} e^{\epsilon k} \right).$$

Proof sketch. The proof is analogous to that of Theorem 5.3. However, without the support hypothesis we cannot guarantee that $\text{supp}(\mathcal{S}(D)) = \text{ext}_D(\mathcal{S}(D'))$, which requires us to find alternative more complex inequalities to bound our results. In particular, we relate a database C^* in the support with $\mathcal{A}_{D, D'}(C^*)$ using the group privacy theorem [17] of approximate DP, i.e.,

$$\mathbb{P}\{\mathcal{M}(C^*) \in A\} \leq e^{\epsilon |C^* \Delta \mathcal{A}_{D, D'}(C^*)|} \mathbb{P}\{\mathcal{M}(\mathcal{A}_{D, D'}(C^*)) \in A\} + \delta \Sigma_{\epsilon}(|C^* \Delta \mathcal{A}_{D, D'}(C^*)|),$$

where $\Sigma_{\epsilon}(l) = \sum_{k=0}^{l-1} e^{\epsilon k}$. We obtain the stated result by substituting this inequality instead of the standard DP one and following the same steps as the proof of Theorem 5.3. The final bound is obtained by grouping all $\mathcal{A}_{D, D'}(C^*)$ of the same value through equalities like

$$\begin{aligned} & \{(C^*, \mathcal{A}_{D, D'}(C^*)) \mid C^* \in \text{supp}(\mathcal{S}(D))\} \\ &= \{(C^*, C) \mid C \in \text{Im}(\mathcal{A}_{D, D'}) \text{ and } C^* \in \mathcal{A}_{D, D'}^{-1}[C]\}. \quad \square \end{aligned}$$

Remark A.6 (On Theorem A.5). We first note that since the result holds for all assignments $\mathcal{A}_{D, D'}$, and thus we can optimize its choice to obtain the smallest $\epsilon_{D, D'}^S$ and $\delta_{D, D'}^S$.

Additionally, Theorem A.5 does generalize Theorem 5.3: If the support condition is verified, then $\text{supp}(\mathcal{S}(D)) = \text{ext}_D(\mathcal{S}(D'))$ for all neighboring databases D and $D' = D_{+y}$. Thus, we can select $\mathcal{A}_{D, D'} = \text{id}$ for all pairs of neighboring databases $D, D' \in \mathbb{D}$, which provides the precise bounds of Theorem 5.3. We expect that this assignment minimizes $\epsilon_{D, D'}^S$ and $\delta_{D, D'}^S$ for the majority of algorithms $\mathcal{S}$. Similarly, we believe that selecting $\mathcal{A}_{D, D'}(C) = C$ for all $C \in \text{supp}(\mathcal{S}(D)) \cap \text{ext}_D(\mathcal{S}(D'))$ is a good choice for the assignment.

A.3 Section 5.3: Distance-Based Probabilistic Suppression

Proof sketch of Theorem 5.4. The proof follows from Theorem 5.3, which can be applied since $\mathcal{S}$ satisfies the support hypothesis (i.e., $\text{supp}(\mathcal{S}(D)) = \{C \mid C \subseteq D\}$ for all $D \in \mathbb{D}$). From Theorem 5.3, we obtain that $\mathcal{M} \circ \mathcal{S}$ satisfy the following inequalities:

$$\begin{aligned} \mathbb{P}\{\mathcal{M}(\mathcal{S}(D)) \in A\} &\leq e^{\epsilon_{D, D'}} \mathbb{P}\{\mathcal{M}(\mathcal{S}(D')) \in A\} + \delta_{D, D'}, \\ \mathbb{P}\{\mathcal{M}(\mathcal{S}(D')) \in A\} &\leq e^{\epsilon_{D', D}} \mathbb{P}\{\mathcal{M}(\mathcal{S}(D)) \in A\} + \delta_{D', D} \end{aligned}$$

for all measurable $A \subseteq \text{Range}(\mathcal{M})$ with (after substituting the expression of $\mathbb{P}\{\mathcal{S}(D) = C\}$)

$$\epsilon_{D, D'}^S = \frac{1}{\text{out}_{D'}(y) + e^{-\epsilon}(1 - \text{out}_{D'}(y))} \cdot \max_{C \subseteq D} \left(\prod_{x \in C} \frac{1 - \text{out}_D(x)}{1 - \text{out}_{D'}(x)} \prod_{x \in D \setminus C} \frac{\text{out}_D(x)}{\text{out}_{D'}(x)} \right),$$

$$\delta_{D, D'}^S = \delta \frac{e^{-\epsilon}(1 - \text{out}_{D'}(y))}{\text{out}_{D'}(y) + e^{-\epsilon}(1 - \text{out}_{D'}(y))},$$

$$\epsilon_{D', D}^S = (\text{out}_{D'}(y) + e^{\epsilon}(1 - \text{out}_{D'}(y))) \cdot \max_{C \subseteq D} \left(\prod_{x \in C} \frac{1 - \text{out}_{D'}(x)}{1 - \text{out}_D(x)} \prod_{x \in D \setminus C} \frac{\text{out}_{D'}(x)}{\text{out}_D(x)} \right),$$

and $\delta_{D', D}^S = \delta(1 - \text{out}_{D'}(y))$.

Parameter δ^S can easily be computed since $\delta_{D, D'}^S \leq \delta_{D', D}^S$ and $\delta^S = \sup_{D \sim D'} \delta(1 - \text{out}_{D'}(y)) = \delta(1 - m)$.

On the other hand, the parameter ϵ^S is much trickier. We rewrite the expressions for $\exp(\epsilon_{D, D'}^S)$ and $\exp(\epsilon_{D', D}^S)$ in terms of $N := |D|$, $T(x, y)$ and $\text{out}_D(x)$ for $x \in D$ by using the identities $\text{out}_{D'}(y) = \frac{1}{N+1}(m + \sum_{x \in D} T(x, y))$ and, for all $x \in D \subseteq D'$,

$$\frac{\text{out}_{D'}(x)}{\text{out}_D(x)} = \frac{N + \frac{T(x, y)}{\text{out}_D(x)}}{N + 1} \quad \text{and} \quad \frac{1 - \text{out}_{D'}(x)}{1 - \text{out}_D(x)} = \frac{N + \frac{1 - T(x, y)}{1 - \text{out}_D(x)}}{N + 1}.$$

The most complex part of the proof is now maximizing over N , $a_{D,D'} := \{T(x,y)\}_{x \in D}$, and $z_D := \{\text{out}_D(x)\}_{x \in D}$ (which corresponds on maximizing over all neighboring D, D'). The key observation for maximizing $\exp(\epsilon_{D',D}^S)$ is that it is convex with respect to z_D , and so the maximum is achieved at its extreme points of its domain. Furthermore, studying then the convexity and concavity of the expression with respect to $a_{D,D'}$ yields an upper bound that depends only on five parameters (including N). We provide a bound over these parameters that is computationally verified (see Remark A.11), obtaining

$$\sup_{\substack{D,D' \in \mathbb{D} \setminus \{\emptyset\}: \\ D'=D+y}} \epsilon_{D',D}^S \leq \max_{p \in [0,1]} \max\{l_1(p), l_2(p)\}.$$

This maximum occurs as $N \rightarrow \infty$ and, intuitively, when D has pN copies of a same element (i.e., at the closest distance from each other) and $(1-p)N$ elements that are the furthest away possible from each other, and y is the furthest away from the former and the closest possible to the latter.

We perform a less complex bound on $\exp(\epsilon_{D,D'}^S)$, since it is usually bounded by that of $\exp(\epsilon_{D',D}^S)$. We obtain

$$\sup_{\substack{D,D' \in \mathbb{D} \setminus \{\emptyset\}: \\ D'=D+y}} \epsilon_{D,D'}^S \leq \max\{e^{l_3}, (e^{-\epsilon} + (1 - e^{-\epsilon})m)^{-1} e^{1-\frac{m}{M}}\},$$

which is also computationally verified (see Remark A.11). By grouping both maximums, we obtain that $\epsilon^S \leq \max_{p \in [0,1]} \max\{l_1(p), l_2(p), l_3\}$, where the remaining term becomes superfluous. The term for $D = \emptyset$, which must be considered separately, is also superfluous. $\square$

Proposition A.7. *Let $\epsilon \geq 0$, and let $m, M \in (0, 1)$ such that $m \leq M$. Let l_1 and l_2 be the functions in $[0, 1]$ as defined in Theorem 5.4.*

Then for $\epsilon \neq 0$ and $m \neq M$, l_1 achieves its maximum over $[0, 1]$ when $p_1 = \min\{1, \max\{V_1, 0\}\}$ with

$$V_1 = \begin{cases} -\frac{1}{3a_1} \left(b_1 + \sqrt[3]{\frac{D_{1,1} + \sqrt{D_{1,1}^2 - 4D_{1,0}^3}}{2}} + \sqrt[3]{\frac{D_{1,1} - \sqrt{D_{1,1}^2 - 4D_{1,0}^3}}{2}} \right) & \text{if } D_{1,1}^2 - 4D_{1,0}^3 > 0 \\ -\frac{1}{3a_1} \left(b_1 + 2\sqrt{D_{1,0}} \cos\left(\frac{1}{3} \arccos\left(\frac{D_{1,1}}{2R_1}\right)\right) \right) & \text{if } D_{1,1}^2 - 4D_{1,0}^3 \leq 0 \end{cases}$$

with

$$\begin{aligned} a_1 &= (e^\epsilon - 1) \frac{M}{m} (M - m)^2, \\ b_1 &= -\frac{M - m}{m} ((m^2 - 4Mm + 2M)(e^\epsilon - 1) + e^\epsilon M), \\ c_1 &= \frac{1 - m}{m} ((e^\epsilon - 1)(2m^2 - 4Mm - m) + (3e^\epsilon - 1)M), \\ d_1 &= -(1 - m) \left((e^\epsilon - 1)(m - 2) + \frac{e^\epsilon}{m} \right), \end{aligned}$$

$$\begin{aligned} D_{1,0} &= b_1^2 - 3a_1 c_1, \\ D_{1,1} &= 2b_1^3 - 9a_1 b_1 c_1 + 27a_1^2 d_1, \\ R_1 &= \sqrt{D_{1,0}^3}; \end{aligned}$$

and l_2 achieves its maximum over $[0, 1]$ when $p_2 = \min\{1, \max\{V_2, 0\}\}$ with

$$V_2 = -\frac{1}{3a_2} \left(b_2 + 2\sqrt{D_{2,0}} \cos\left(\frac{1}{3} \arccos\left(\frac{D_{2,1}}{2R_2}\right)\right) \right)$$

with

$$\begin{aligned} a_2 &= \frac{e^\epsilon - (e^\epsilon - 1)m}{m}, \\ b_2 &= -\frac{6e^\epsilon - (e^\epsilon - 1)(M + 5m)}{m}, \\ c_2 &= \frac{m((e^\epsilon - 1)(m + 9M - 9) - e^\epsilon) + 4M((e^\epsilon - 1)M - 4e^\epsilon + 1) + 12e^\epsilon}{(1 - M)m}, \\ d_2 &= -(2e^\epsilon - (e^\epsilon - 1)(M + m)) \left(\frac{4 - 4M - m}{(1 - M)m} \right) + 2(e^\epsilon - 1), \\ D_{2,0} &= b_2^2 - 3a_2 c_2, \\ D_{2,1} &= 2b_2^3 - 9a_2 b_2 c_2 + 27a_2^2 d_2, \\ R_2 &= \sqrt{D_{2,0}^3}. \end{aligned}$$

For $\epsilon = 0$ and $m \neq M$, l_1 achieves its maximum over $[m, M]$ when $p_1 = \min\{1, \max\{V_1, 0\}\}$ with

$$V_1 = \frac{1 - m}{M - m} - \frac{\sqrt{Mm(1 - m)(1 - M)}}{M(M - m)},$$

and l_2 achieves its maximum over $[0, 1]$ when $p_2 = \min\{1, \max\{V_2, 0\}\}$ with

$$V_2 = 2 - \frac{\sqrt{m(1 - M)}}{1 - M}.$$

For $m = M$, l_1 and l_2 are constantly $\ln(e^\epsilon - (e^\epsilon - 1)m)$.

Proof. We find the maximum of l_i (for $i \in [2]$) by studying derivatives. For simplicity, we compute those of $L_i := \ln \circ l_i$ since the maximums of l_i and L_i are the equivalent because the function $\ln$ is strictly increasing. Thus, we have that

$$\frac{\partial L_i}{\partial p}(p) = -(M - m) \frac{a_i p^3 + b_i p^2 + c_i p + d_i}{(p - A_i)^{n_i}} K_{m,M}(p),$$

where a_i, b_i, c_i , and d_i are the constants in the statement, $K_{m,M}(p) > 0$ for all $p \in \mathbb{R}$, $A_i \geq 1$ and $n_i \in \mathbb{N}$. The extreme points of L_i are thus the roots of $a_i p^3 + b_i p^2 + c_i p + d_i = 0$ that are not A_i , the asymptote of L_i . The roots of this cubic polynomial are found using the general cubic formula (note that the degenerate cases $\epsilon = 0$ and $m = M$ need to be considered apart). Studying these roots reveals that the smallest real root V_i (as shown in the statement) is the only critical point that can be in the interval $[0, 1]$, and that L_i is increasing in $(-\infty, V_i)$ and decreasing in (V_i, α_i) where $1 < \alpha_i \leq A_i$. Thus, L_i achieves its maximum at $p_i := \min\{1, \max\{V_i, 0\}\}$. $\square$

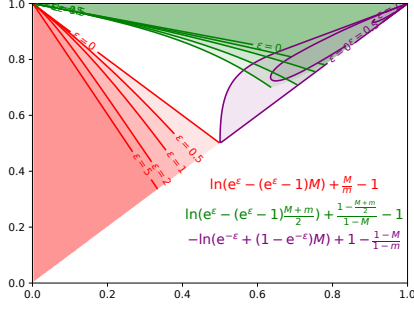


Figure 4: The colored areas show the values of m (x -axis) and M (y -axis) such that $\epsilon^S = \epsilon^S(\epsilon, m, M)$ simplifies, which depends on the given ϵ . The expression of ϵ^S simplifies for the values m and M within a colored region to the expression of the same color. These colored areas are also where the privacy parameters are tight with respect to Theorem 5.3.

Remark A.8 (On the tightness of the privacy parameters in Theorem 5.4). The bounds we provide are tight in the colored areas in Figure 4 with respect to Theorem 5.3 (see long version), but we are not able to show if they are tight in general. The pair of neighboring databases D and D_{+y} that achieves these “tight” bounds corresponds to having all elements of D be the closest possible records (i.e., all copies of the same element x) and y be the furthest record from x (red area); and having all elements of D be the furthest away from each other and y be the closest to all of them (green and purple areas; possible under certain metrics). We believe that the values that achieve the tight bounds in the uncolored region would correspond to a combination of these cases, i.e., a portion of records are as close as possible, while the rest are as far away as possible. In this case, the portion would be determined by the value p that achieves the maximum.

Proposition A.9. *Let S be a suppression algorithm with domain $\mathbb{D}$ and let $\epsilon_{\text{tight}}^S$ and δ_{tight}^S be the tight privacy parameters of $\mathcal{M} \circ S$ over all (ϵ, δ) -DP mechanisms $\mathcal{M}$ with domain $\mathbb{D}$ ($\delta \leq 1$), i.e., for all (ϵ, δ) -DP mechanism $\mathcal{M}$, we have that $\mathcal{M} \circ S$ is $(\epsilon_{\text{tight}}^S, \delta_{\text{tight}}^S)$ -DP and there does not exist $\epsilon'' \leq \epsilon_{\text{tight}}^S$ and $\delta'' \leq \delta_{\text{tight}}^S$ (not both equal) such that $\mathcal{M} \circ S$ is (ϵ'', δ'') -DP for all (ϵ, δ) -DP mechanisms $\mathcal{M}$. Then, $\epsilon_{\text{tight}}^S \geq \ln(1 + (e^\epsilon - 1)p)$ and $\delta_{\text{tight}}^S \geq \delta p$ with $p = \sup_{D' \in \mathbb{D}} \sup_{y \in D'} \mathbb{P}\{y \in S(D')\}$. The right terms of the inequalities are precisely the privacy parameters of the uniform Poisson sampling with sampling rate p .*

Proof sketch. It suffices to show that there exists an (ϵ, δ) -DP mechanism $\mathcal{M}$ such that $\mathcal{M} \circ S$ is tightly (ϵ'', δ'') -DP with $\epsilon'' = \ln(1 + (e^\epsilon - 1)p)$ and $\delta'' = \delta p$.

We consider the mechanism $\mathcal{M}: \mathbb{D} \rightarrow \{0, 1\}$ such that

$$\mathbb{P}\{\mathcal{M}(C) = 0\} = \begin{cases} \frac{1-\delta}{1+e^\epsilon} & \text{if } y \notin C, \\ \frac{e^\epsilon + \delta}{1+e^\epsilon} & \text{if } y \in C, \end{cases}$$

and $\mathbb{P}\{\mathcal{M}(C) = 1\} = 1 - \mathbb{P}\{\mathcal{M}(C) = 0\}$. This mechanism is well-defined and satisfies (ϵ, δ) -DP tightly.

Now, by fixing D, D' be two neighboring databases such that $D' = D_{+y}$, we can compute the probability distributions of $\mathcal{M}(S(D))$ and $\mathcal{M}(S(D'))$ obtaining

$$\mathbb{P}\{\mathcal{M}(S(D)) \in A\} \leq e^{\epsilon_{D,D'}} \mathbb{P}\{\mathcal{M}(S(D')) \in A\} + \delta_{D,D'}$$

and

$$\mathbb{P}\{\mathcal{M}(S(D')) \in A\} \leq e^{\epsilon_{D',D}} \mathbb{P}\{\mathcal{M}(S(D)) \in A\} + \delta_{D',D}$$

with $\epsilon_{D,D'} = \epsilon_{D',D} = \ln(e^\epsilon q + (1-q))$ and $\delta_{D,D'} = \delta_{D',D} = \delta q$ for all subsets $A \subseteq \{0, 1\}$, such that an inequality is tight for at least one subset (in this case, the second inequality for $A = \{0\}$). Taking the supremum over all neighboring $D, D' \in \mathbb{D}$, we obtain that $\mathcal{M} \circ S$ is tightly (ϵ'', δ'') -DP, which, by construction, verifies $\epsilon'' \leq \epsilon_{\text{tight}}^S$ and $\delta'' \leq \delta_{\text{tight}}^S$, following the notation of the statement. $\square$

Corollary A.10. *The outlier-score suppression with parameters m and M cannot provide a greater privacy amplification than the uniform Poisson sampling with sampling rate $1 - m$.*

Proof. Follows directly from Proposition A.9 using that $\sup_{D' \in \mathbb{D}} \sup_{y \in D'} \mathbb{P}\{y \in S(D')\} = 1 - m$. $\square$

Remark A.11 (Numerical computation). We resort to numerical computations in the last steps of the calculation of the privacy parameter ϵ^S of Theorem 5.4. More precisely, we require computational power to compute two maximizations with respect to two and five parameters, respectively. One parameter in each optimization, $N \in \mathbb{N}$, is unbounded, and thus we check the maximum for up to $N = 10^9$, which we deem more than enough since the function convergences quickly when $N \rightarrow \infty$. All other parameters are bounded.

We perform the optimizations using `find_local_maximum` of the `sage` package and `differential_evolution` of the `script.optimize` package in Python.

Since the domains of $\epsilon \in [0, \infty)$, $m \in (0, 1)$ and $M \in [m, 1)$ are continuous and the domain of ϵ is furthermore unbounded, we cannot run exhaustively for all values of these parameters. For this work, we decide to run the experiment for every reasonable value with decent level of granularity: m and M are run for every in $\{0.01, 0.02, \dots, 0.98, 0.99\}$ (with $m \leq M$) and ϵ is run for every value in $\{0, 0.01, 0.02, \dots, 1.98, 1.99, 2\}$, $\{2.1, 2.2, \dots, 9.9, 10\}$, and $\{11, 12, \dots, 99, 100\}$. We believe that our choices for ϵ are representative of the values of ϵ deemed to be acceptable in the literature (small values smaller than 2 or 10 [17]).

Our experiments show that the hypothesized value is the correct bound up to an error of $2 \cdot 10^{-7}$. In addition, since the evaluated functions are continuous and smooth with respect to ϵ, m and M , we conjecture that this is the real bound for all $\epsilon \geq 0$ and $m, M \in (0, 1)$ such that $m \leq M$.

We note that the full proofs of our results can be found in Section B of the long version (arXiv:2601.05180).