

PrivRM: A Framework for Range Mean Estimation under Local Differential Privacy

LIANTONG YU, The Hong Kong Polytechnic University, Hong Kong SAR, China

QINGQING YE*, The Hong Kong Polytechnic University, Hong Kong SAR, China

RONG DU, The Hong Kong Polytechnic University, Hong Kong SAR, China

The increasing collection and analysis of personal data driven by digital technologies has raised concerns about individual privacy. Local Differential Privacy (LDP) has emerged as a promising solution to provide rigorous privacy guarantee for users, without relying on a trusted data collector. In the context of LDP, range mean estimation over numerical values is an important yet challenging problem. Simply applying existing work may introduce overly large noise sensitivity, since all of them focus on statistical tasks (e.g., mean or distribution) across the entire domain. In this paper, we propose a novel framework for Private Range Mean (PrivRM) estimation under LDP. Two implementations of the framework, namely *PrivRM^L* and *PrivRM^{*}*, are developed, which are adaptable to all existing numerical value perturbation mechanisms. As an optimization of the framework, we also propose a distribution-aware Adaptive Adjustment (AA) strategy to dynamically confine the perturbation space for skewed data distributions. Extensive experimental results show that under the same privacy guarantee and query range, our framework *PrivRM* significantly improve over existing solutions.

CCS Concepts: • **Security and privacy** → **Data anonymization and sanitization**.

Additional Key Words and Phrases: Local Differential Privacy; Range Mean Estimation.

ACM Reference Format:

Liantong Yu, Qingqing Ye, and Rong Du. 2025. *PrivRM: A Framework for Range Mean Estimation under Local Differential Privacy*. *Proc. ACM Manag. Data* 3, 3 (SIGMOD), Article 205 (June 2025), 26 pages. <https://doi.org/10.1145/3725414>

1 Introduction

With the increasing popularity of big data analytics, data collectors are becoming more interested in collecting and analyzing usage data from their customers in order to improve their services. Nevertheless, the practice of data sharing carries the risk of data leakage, which can affect not only users but also third parties susceptible to both internal and external data leakage. To safeguard personal data, Local Differential Privacy (LDP) [9, 16, 18] has been proposed to provide a robust privacy-preserving method for various data analysis tasks. By enabling users to perturb their data locally before sharing, LDP ensures that sensitive information remains confidential even in the presence of an untrusted data collector. This paradigm has far-reaching implications for the security and privacy of data in real-world applications, such as Apple [29], Google [20] and Microsoft [11].

In the context of LDP, mean estimation over numerical values is a fundamental query and attracts much attention from researchers [17, 19, 30]. In the literature, all the existing work focus on mean

*Corresponding author.

Authors' Contact Information: Liantong Yu, The Hong Kong Polytechnic University, Hong Kong SAR, China, liantong2001.yu@connect.polyu.hk; Qingqing Ye, The Hong Kong Polytechnic University, Hong Kong SAR, China, qqing.ye@polyu.edu.hk; Rong Du, The Hong Kong Polytechnic University, Hong Kong SAR, China, roong.du@connect.polyu.hk.



This work is licensed under a Creative Commons Attribution 4.0 International License.

© 2025 Copyright held by the owner/author(s).

ACM 2836-6573/2025/6-ART205

<https://doi.org/10.1145/3725414>

estimation over the entire domain. However, in practical applications, the query range is usually specific. Below are two examples.

- **Employee Salary Analysis.** Consider a substantial company with million of employees, the Human Resources (HR) department may seek to ascertain the average salary levels of employees falling within a specific salary range (e.g., annual salaries between \$50, 000 and \$100, 000) to facilitate compensation adjustments and budget planning.
- **House Price Analysis.** In extensive real estate datasets, the transaction price of each property is recorded. Real estate agents or potential purchasers may seek to know the transaction selling price of houses within a specific price range (e.g., from \$300, 000 to \$500, 000) to identify market trends and make informed purchasing decisions.

The two examples can be systematically formulated as a problem of range mean estimation. Without loss of generality, we assume each user possesses a value drawn from the domain $\mathcal{D}$. Given a specific range $[l, r] \subseteq \mathcal{D}$, the objective is to return the mean of values which fall within the range, while ensuring LDP guarantees.

In the literature, several studies focus on LDP-enabled numerical value perturbation, including the Laplace Mechanism (LM) [19], Stochastic Rounding (SR) [18], Piecewise Mechanism (PM), Hybrid Mechanism (HM) [30], and Square Wave (SW) mechanism [26]. These mechanisms are specifically designed for mean or distribution estimation. However, they all consider estimation over the entire domain, leaving a gap in range-specific mean estimation. Notably, directly applying these mechanisms for mean estimation over a specific range results in significant utility loss, as their noise sensitivity typically covers the entire domain. This issue is particularly pronounced when the query range $[l, r]$ is substantially smaller than the entire domain.

Intuitively, a range mean depends solely on the values within that range and is irrelevant to any values outside it. Therefore, an important way to improve the utility of range mean estimation is to concentrate on the query range to reduce the noise sensitivity. However, under LDP, a key challenge is how to perturb out-of-range values while eliminating their influence on the aggregation. Based on this idea, we propose a framework for locally differentially Private Range Mean (*PrivRM*) estimation. This framework reduces the sensitivity from the entire domain to the query range by dividing the task into two phases: estimating range count in Phase 1 and the range sum in Phase 2. We implement the framework *PrivRM* in two ways, namely *PrivRM^l* and *PrivRM^{*}*. The first one eliminates the influence of the out-of-range values on the mean by employing randomization techniques to the input domains. Besides, we discover that while employing randomization techniques to the output domains, we can leverage more privacy budget across both phases while still providing the same LDP guarantee. Based on this, we present an optimized implementation, *PrivRM^{*}*, which exclusively employs in-range and out-of-range values for range mean estimation. Our observations indicate that compared to *PrivRM^l*, *PrivRM^{*}* consistently introduces less perturbation noise while ensuring the same LDP guarantee, resulting in superior data utility.

To further optimize the performance, we design a distribution-aware strategy that can be integrated into our framework. When estimating the range mean over a skewed distribution, where values are concentrated in an area much smaller than the query range, we can confine the perturbation space to a densely populated region to further reduce the noise sensitivity. Towards this idea, we introduce an Adaptive Adjustment (AA) strategy. AA strategy employs a binary search within the query range to dynamically confine the perturbation space, ensuring that the perturbation mechanism is only applied to the dense region, thus preserving the utility of data while ensuring LDP guarantee.

Overall, the main contributions of this paper are as follows.

- To our knowledge, this is the first work to explore the problem of range mean estimation under LDP. We propose a novel framework *PrivRM*, which is capable of integrating all the existing numerical mechanisms.
- We implement the framework in two ways, namely *PrivRM^l* and *PrivRM^{*}*, based on which a black-box guideline is provided to offer the best suggestion.
- We advance the framework with our Adaptive Adjustment (AA) strategy to optimize the utility, especially for skewed data distributions.
- Extensive experimental results on real-world datasets validate the effectiveness of *PrivRM* framework in supporting range mean query.

Roadmap. Section 2 introduces the preliminaries on LDP. Section 3 presents the problem definition and the existing solutions. Sections 4 and 5 introduce the *PrivRM* framework and its implementations. A framework optimization is proposed in Section 6. Section 7 extends the regular range mean estimation to multidimensional and multiple queries cases, and discusses the relationship between range mean estimation and general range query. Experimental results are presented in Section 8. Finally, we discuss related work in Section 9 and conclude the paper in Section 10.

2 Preliminaries

2.1 Local Differential Privacy

Differential privacy (DP) works in both centralized and local settings. Centralized DP requires the data curator to be fully trusted to collect all data [19], while local DP does not rely on this assumption [16]. In the local setting, each user locally perturbs her data before reporting them to an untrusted data collector, which makes it more secure and practical in real-world applications. The formal definition is as follows.

DEFINITION 1. (Local Differential Privacy, LDP) *A randomized algorithm $\mathcal{A}$ satisfies ϵ -LDP if for any two input records w and w' , and any output w^* of $\mathcal{A}$, the following inequality holds.*

$$\Pr[\mathcal{A}(w) = w^*] \leq e^\epsilon \times \Pr[\mathcal{A}(w') = w^*]. \quad (1)$$

In Equation 1, ϵ is called the privacy budget, which controls the deniability of a randomized algorithm taking w or w' as its input. As with centralized DP, LDP also has the property of sequential composition as below, which guarantees the overall privacy for a sequence of randomized algorithms.

THEOREM 1. (Sequential Composition) *Given c randomized algorithms $\mathcal{A}_i (1 \leq i \leq c)$, each providing ϵ_i -local differential privacy. Then the sequence of algorithms $\mathcal{A}_i (1 \leq i \leq c)$ collectively provides $(\sum \epsilon_i)$ -local differential privacy.*

2.2 Randomized Response

The technique of randomized response (RR) [33] serves as a paradigm to ensure ϵ -LDP, which has been widely adopted in LDP perturbation mechanisms. Specifically, RR enables respondents to answer a sensitive binary question while maintaining plausible deniability. In essence, each user reports a genuine answer with probability p and provides false answer with probability $1 - p$. To comply with ϵ -LDP, RR sets the probability $p = \frac{e^\epsilon}{1+e^\epsilon}$, so that $\frac{p}{1-p}$ is bounded by e^ϵ .

Various LDP mechanisms have been developed in accordance with RR, depending on whether the data type is categorical or numerical. This work primarily focuses on the LDP perturbation mechanisms for numerical values, which are detailed in the following section.

2.3 LDP Mechanisms for Numerical Values

We introduce five typical numerical values perturbation (NVP) mechanisms, namely Laplace Mechanism (LM) [19], Stochastic Rounding (SR) [18], Piecewise Mechanism (PM) [30], Hybrid Mechanism (HM) [30] and Square Wave Mechanism (SW) [26]. Without loss of generality, we assume the domain of an input value is $[-1, 1]$ in the sequel. Note that LM is an unbounded mechanism generating random noise from $[-\infty, +\infty]$, while the others are bounded mechanisms generating noise from a specific range.

2.3.1 Laplace Mechanism (LM). LM [19] adds noise randomly drawn from Laplace distribution to data, ensuring that individual privacy is preserved while allowing statistical analysis. Given an input $t_i \in [-1, 1]$ and privacy budget ϵ , the output of LM is defined as $\hat{t}_i = t_i + \text{Lap}(\frac{2}{\epsilon})$, where $\text{Lap}(\frac{2}{\epsilon})$ is a random variable that conforms to the Laplace distribution with the probability density function $pdf(x) = \frac{\epsilon}{4} \exp\left(-\frac{\epsilon|x|}{2}\right)$.

2.3.2 Stochastic Rounding (SR). The essence of SR [18] is the Bernoulli Distribution, so the output is bounded and discrete. Given an input value $t_i \in [-1, 1]$ and privacy budget ϵ , the output follows a Bernoulli Distribution, taking the value from $\{-C_{SR}(\epsilon), C_{SR}(\epsilon)\}$, where $C_{SR}(\epsilon) = \frac{e^\epsilon + 1}{e^\epsilon - 1}$:

$$\Pr[\hat{t}_i = y \mid t_i] = \begin{cases} \frac{e^\epsilon - 1}{2e^\epsilon + 2} \cdot t_i + \frac{1}{2}, & \text{if } y = C_{SR}(\epsilon), \\ -\frac{e^\epsilon - 1}{2e^\epsilon + 2} \cdot t_i + \frac{1}{2}, & \text{if } y = -C_{SR}(\epsilon). \end{cases}$$

2.3.3 Piecewise Mechanism (PM). Compared to SR, the output of PM is continuous [30]. Given an input value $t_i \in [-1, 1]$, PM outputs a value $\hat{t}_i \in [-C_{PM}(\epsilon), C_{PM}(\epsilon)]$, where $C_{PM}(\epsilon) = \frac{e^{\epsilon/2} + 1}{e^{\epsilon/2} - 1}$. Particularly, PM perturbs the data using a piecewise probability density function that consists of an interval and other intervals:

$$\Pr[\hat{t}_i = y \mid t_i] = \begin{cases} p = \frac{(e^\epsilon - e^{\epsilon/2})}{2(e^{\epsilon/2} + 1)}, & \text{if } y \in [l(t_i), r(t_i)], \\ q = \frac{e^{\epsilon/2} - 1}{2(e^{\epsilon/2} + e^\epsilon)}, & \text{otherwise,} \end{cases}$$

where $l(t_i) = \frac{e^{\epsilon/2}t_i - 1}{e^{\epsilon/2} - 1}$ and $r(t_i) = \frac{e^{\epsilon/2}t_i + 1}{e^{\epsilon/2} - 1}$.

2.3.4 Hybrid Mechanism (HM). When the data distribution is concentrated in the middle of the input domain, PM performs better than SR; otherwise, SR shows superior performance [30]. Here, HM combines both approaches by achieving lower worst-case variance of SR and PM [30]. Specifically, HM operates by invoking PM with probability α_{HM} while invoking SR with probability $1 - \alpha_{HM}$, where

$$\alpha_{HM} = \begin{cases} 1 - e^{-\epsilon/2}, & \text{if } \epsilon > 0.61, \\ 0, & \text{if } \epsilon \leq 0.61. \end{cases}$$

It has been proven that these four mechanisms are unbiased, i.e., $\mathbb{E}(\hat{t}_i) = t_i$.

2.3.5 Square Wave Mechanism (SW). Originally, SW [26] perturbs a value $t_i \in [0, 1]$ into a sanitized version $\hat{t}_i \in [-b, 1 + b]$ via Equation 2

$$\Pr[\hat{t}_i = y \mid t_i] = \begin{cases} p = \frac{e^\epsilon}{2be^\epsilon + 1}, & \text{if } |t_i - y| \leq b, \\ q = \frac{1}{2be^\epsilon + 1}, & \text{otherwise,} \end{cases} \quad (2)$$

where $b = \frac{e^\epsilon - e^{\epsilon/2}}{2e^\epsilon(e^\epsilon - 1 - e^{\epsilon/2})}$. To handle the unified domain $[-1, 1]$, we can normalize t_i into $[-1, 1]$, perturb it via Equation 2, and then denormalize the perturbed value. By merging normalization

Table 1. Notation

Symbol	Description
$\mathcal{U}$	the set of users
n	the total number of users, $n = \mathcal{U} $
u_i	the i -th user
$\mathcal{T}$	the set of values
t_i	the value possessed by u_i
$n_{(l,r)}$	the number of users with values in range $[l, r]$
$s_{(l,r)}$	the summation of values in range $[l, r]$
$m_{(l,r)}$	the mean of values in range $[l, r]$
$\mathbf{d}_{(l,r)}$	the distribution of values in range $[l, r]$
$\mathbb{M}_{\mathcal{A}}(\cdot)$	the output domain of NVP mechanism $\mathcal{A}$

and denormalization with Equation 2, the perturbation of SW can be represented as

$$\Pr [\hat{t}_i = y \mid t_i] = \begin{cases} p = \frac{e^\epsilon}{2(2be^\epsilon + 1)}, & \text{if } |t_i - y| \leq 2b, \\ q = \frac{1}{2(2be^\epsilon + 1)}, & \text{otherwise,} \end{cases} \quad (3)$$

which takes a value $t_i \in [-1, 1]$ as input and outputs $\hat{t}_i \in [-1 - 2b, 1 + 2b]$.

Note that, unlike LM, SR, PM and HM which ensure unbiasedness, SW is not an unbiased estimator [15]. While this is to reduce the estimation variance, it can only cope with relatively symmetric distribution well. For the others, especially those skew distributions, SW may cause much estimation deviation from the ground truth. Fortunately, this bias can be corrected by applying a correction processing as

$$\hat{t}_i = \frac{y}{4b(p - q)} \quad (4)$$

The proof of unbiasedness of SW is presented in the full version [3].

After the perturbation (Equation 2 or 3), in order to reconstruct the histogram of the data distribution, the aggregator needs to divide both the input domain and the output domain into several bins, typically 1024 bins. And then collector reconstructs the distribution by using an Expectation Maximization algorithm (EM).

3 Problem Definition and Existing solutions

In this section, we first formulate the problem of range mean estimation under LDP, and then present two baseline solutions that are directly adapted from existing works.

3.1 Problem Definition

We consider a system setting where there are a set of n users $\mathcal{U} = \{u_1, u_2, \dots, u_n\}$, and a collector. Each user u_i holds a numerical value $t_i \in [-1, 1]$, where $[-1, 1]$ is the value domain. Given a specific range $[l, r] \subseteq [-1, 1]$, the collector's objective is to calculate the mean $m_{(l,r)}$ of users values which fall within $[l, r]$, while ensuring the data privacy of all users. Formally,

$$m_{(l,r)} = \frac{s_{(l,r)}}{n_{(l,r)}} = \frac{\sum_{i=1}^n t_i \cdot \mathbb{1}_{[l,r]}(t_i)}{\sum_{i=1}^n \mathbb{1}_{[l,r]}(t_i)}, \quad (5)$$

where $\mathbb{1}_{[l,r]}(t_i)$ is an indicator function that returns 1 if $t_i \in [l, r]$ and 0 otherwise.

3.2 Existing Solutions

To estimate the range mean as shown in Equation 5, there are two baseline solutions adapted from existing works.

Solution 1: Direct Estimation. Each user perturbs her data t_i and reports a sanitized value $\hat{t}_i$ by a numerical value perturbation mechanism. Then the collector derives the mean $\hat{m}_{(l,r)}$ by calculating the sum and count of values within the range $[l, r]$. Specifically,

$$\hat{m}_{(l,r)} = \frac{\sum_{i=1}^n \hat{t}_i \cdot \mathbb{1}_{[l,r]}(\hat{t}_i)}{\sum_{i=1}^n \mathbb{1}_{[l,r]}(\hat{t}_i)}.$$

Note that only those NVP mechanisms whose output domains are continuous can be adopted, e.g., LM [19], PM [30] and SW [26].

Solution 2: Distribution-based Estimation. Another idea is to approximate the distribution of user values, based on which we can then estimate the mean of a specific range. SW mechanism [26] supports such a solution. Assume the approximated distribution by SW is denoted by $\hat{\mathbf{d}}$. Given a specific range $[l, r]$, the solution works as follows. First, the bins of the specified range, represented by the set $\{t_{(l,r)}^1, \dots, t_{(l,r)}^k\}$, are filtered, along with their corresponding probability densities, represented by the set $\{\hat{d}_{(l,r)}^1, \dots, \hat{d}_{(l,r)}^k\}$. The mean of in-range data can be calculated by dividing the corresponding sum by the count of values. Formally,

$$\hat{m}_{(l,r)} = \frac{n \sum_{i=1}^k t_{(l,r)}^i \hat{d}_{(l,r)}^i}{n \sum_{i=1}^k \hat{d}_{(l,r)}^i} = \frac{\sum_{i=1}^k t_{(l,r)}^i \hat{d}_{(l,r)}^i}{\sum_{i=1}^k \hat{d}_{(l,r)}^i}.$$

Indeed, a range mean query primarily focuses on the values within a specified range. However, to prevent the disclosure of other values, both Solutions 1 and 2 treat all values equally, and consider the overall domain length 2 as the sensitivity when adding noise to a value. These broad-brush solutions do not differentiate between values within and outside the specified range, thus diverting attention from the specified range, which may result in significant utility loss.

More specifically, in Solution 1, the estimated mean is unbiased only if the effects of the data entering and exiting the range after the perturbation are cancelled out. As for Solution 2, SW needs to approximate the value distribution on the whole domain, dividing it into 1024 blocks. When given a small query range, such as 1/1024 of the domain, SW will treat that region as a single block, ignoring the distribution within that region and assuming a uniform distribution by default. This ultimately leads to large non-uniform estimation error.

In a nut shell, for range mean query, it is beneficial to have a more fine-grained solution that focuses on the values within the specified range. This could potentially lead to a reduction in utility loss.

4 PrivRM: Private Range Mean Estimation

In this section, we present a general framework for locally differentially Private Range Mean estimation (*PrivRM*), which differentiates values within and outside the specified range to enhance the estimation accuracy. In what follows, we will present an overview of *PrivRM* in Section 4.1 and then introduce the first implementation in Sections 4.2.

4.1 Framework Overview

To concentrate on the target range rather than the entire domain, a sensible strategy is to confine the perturbed values to be within the range, while the impact of values initially outside this range can be eliminated from the aggregation. Inspired by this, we introduce the *PrivRM* framework for

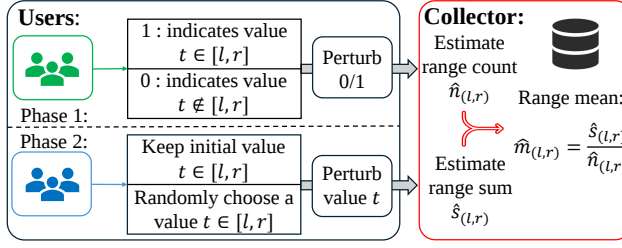


Fig. 1. The overview of PrivRM framework

estimating the range mean, and Figure 1 shows an overview of it. Specifically, PrivRM first employs the RR mechanism to perturb whether each user's value falls within the range $[l, r]$, estimating the count of values within this range (a.k.a., range count), denoted by $\hat{n}_{(l,r)}$. Subsequently, it uses an NVP mechanism (see Section 2.3) to perturb the values, estimating the sum of values that are initially within the range (a.k.a., range sum). The range mean is then calculated by dividing this range sum by the range count. However, a challenge remains: how to perturb the values outside the range $[l, r]$ in a way that satisfies ϵ -LDP and produces an unbiased sum $\hat{s}_{(l,r)}$, free from the influence of those outside values.

Formally, we can calculate the range sum $\hat{s}_{(l,r)}$ as

$$\hat{s}_{(l,r)} = \sum_{i=1}^n \hat{t}_i - (n - \hat{n}_{(l,r)})\mathbb{E}[\hat{t}_{out}],$$

where $\mathbb{E}[\hat{t}_{out}]$ is the expectation of the perturbed values initially outside the range $[l, r]$. Then we can estimate the range mean as

$$\hat{m}_{(l,r)} = \frac{\hat{s}_{(l,r)}}{\hat{n}_{(l,r)}} = \frac{\sum_{i=1}^n \hat{t}_i - (n - \hat{n}_{(l,r)})\mathbb{E}[\hat{t}_{out}]}{\hat{n}_{(l,r)}} \quad (6)$$

THEOREM 2. In PrivRM, with an unbiased NVP mechanism, the estimated sum is also unbiased, i.e., $\mathbb{E}[\hat{s}_{(l,r)}] = s_{(l,r)}$. And the variance of range mean is

$$\text{Var}(\hat{m}_{(l,r)}) \approx \frac{\mathbb{E}^2[X]}{\mathbb{E}^2[Y]} \left(\frac{\text{Var}(X)}{\mathbb{E}^2[X]} + \frac{\text{Var}(Y)}{\mathbb{E}^2[Y]} \right) \quad (7)$$

where $X = \sum_{i=1}^n \hat{t}_i - n\mathbb{E}[\hat{t}_{out}]$ and $Y = \hat{n}_{(l,r)}$.

PROOF. The proof appears in the full version [3]. □

According to Equation 6, range mean estimation also relies on $\mathbb{E}[\hat{t}_{out}]$. To facilitate its estimation, we may apply input or output domain randomization to the values outside the range $[l, r]$. Firstly, we design an implementation based on input randomization, namely PrivRM^I.

4.2 PrivRM^I: Input Domain Randomization

We first introduce input-randomization-based implementation of the framework, namely PrivRM^I. The idea is to randomize the values initially outside the query range $[l, r]$ into random values within this range. Therefore, the randomized values, which become the inputs of the NVP mechanism, satisfy a uniform distribution over the range $[l, r]$. After applying an NVP mechanism to these values which are initially outside the range, the expectation of their perturbed values is kept fixed, i.e., $\mathbb{E}[\hat{t}_{out}] = \frac{l+r}{2}$.

Algorithm 1 shows the pseudo-code of PrivRM^I. It involves two phases, each of which consumes half of the given privacy budget ϵ . In Phase 1, each user employ RR to report whether their value is within the range $[l, r]$ (Lines 3-6). Then in Line 7, the collector derives the range count $\hat{n}_{(l,r)}$ by

Algorithm 1: Workflow of PrivRM^I **Input:** A set of user values $\{t_1, t_2, \dots, t_n\}$, query range $[l, r]$, NVP mechanism $\mathcal{A}(\cdot)$, privacy budget ϵ **Output:** The estimated range mean $\hat{m}_{(l,r)}$

```

1 for  $i = 1$  to  $n$  do // Phase 1
2   Set  $p = \frac{e^{\epsilon/2}}{1+e^{\epsilon/2}}$ 
3   if  $t_i \in [l, r]$  then
4     Report bit 1 (resp. 0) with probability  $p$  (resp.  $1 - p$ )
5   else
6     Report bit 0 (resp. 1) with probability  $p$  (resp.  $1 - p$ )
7 Collector derives the range count  $\hat{n}_{(l,r)} = \frac{n(p-1)+\#('1')}{2p-1}$ 
8 for  $i = 1$  to  $n$  do // Phase 2
9   if  $t_i \in [l, r]$  then
10    Normalization:  $t'_i = \frac{2(t_i-l)}{r-l} - 1$ 
11    Perturbation:  $y_i = \mathcal{A}(t'_i, \epsilon/2)$ 
12    Denormalization:  $\hat{t}_i = \frac{(y_i+1)(r-l)}{2} + l$ 
13  else
14    Randomly draw a value  $t'_i \in [-1, 1]$ 
15    Perturbation:  $y_i = \mathcal{A}(t'_i, \epsilon/2)$ 
16    Denormalization:  $\hat{t}_i = \frac{(y_i+1)(r-l)}{2} + l$ 
17 Calculate  $\hat{m}_{(l,r)}$  by Equation 6, where  $\mathbb{E}[\hat{t}_{out}] = \frac{l+r}{2}$ 
18 return  $\hat{m}_{(l,r)}$ 

```

taking noise calibration into account. In Phase 2, we focus on estimating the range sum. Since the value domain of the NVP mechanisms is unified as $[-1, 1]$, before perturbing a value $t_i \in [l, r]$, we will first normalize it into $t'_i \in [-1, 1]$ (Line 10). The perturbation is then applied to t'_i with the remaining privacy budget $\epsilon/2$ (Line 11). Subsequently, we will also denormalize the perturbed value y_i (Line 12). For a value outside the range $[l, r]$, we draw a random value from the input domain, i.e., $t'_i \in [-1, 1]$, perturb it and then apply denormalization (Lines 14-16). Here since the random value t'_i is drawn from the input domain $[-1, 1]$, which is often different from the output domain of an NVP mechanism, t'_i will be then perturbed to ensure the indistinguishability of output values initially within and outside the range. Ultimately, we can estimate the range mean via Equation 6 (Line 17).

In PrivRM^I , all the involved perturbation mechanisms (including RR, LM, SR, PM, HM and Unbiased SW) are all unbiased. Thus we have

$$\mathbb{E}[X] = s_{(l,r)} - n_{(l,r)} \left(\frac{l+r}{2} \right), \mathbb{E}[Y] = n_{(l,r)}.$$

As for the variance,

$$\text{Var}(X) = n \cdot \text{Var}(\epsilon/2), \text{Var}(Y) = \frac{ne^{\epsilon/2}}{(e^{\epsilon/2} - 1)^2},$$

where $\text{Var}(\epsilon/2)$ is the variance of a NVP mechanism with privacy budget $\epsilon/2$. We summarize the variance of each NVP mechanism in the full version [3]. Subsequently, the variance of PrivRM^I can be calculated through Equation 7.

5 PrivRM*: An Optimized Implementation

In PrivRM^l , both randomization and perturbation are applied to the values initially outside the range $[l, r]$, since the randomization domain $[-1, 1]$ may be different from the output domain of an NVP mechanism $\mathcal{A}(\cdot)$. Inspired by this, we may alternatively draw a random value from the output domain $\mathbb{M}_{\mathcal{A}}(\cdot)$, so that we can remove the perturbation step in Phase 2. Besides, PrivRM^l completely separates the process into two distinct phases, allocating half of the privacy budget to each phase. This division inevitably leads to an increase in the scale of noise introduced. In this section, we present an optimized implementation based on output randomization, PrivRM^* , which allows us to leverage more privacy budget across both phases while still providing LDP guarantee, thus achieving better utility.

5.1 Design Rationale of PrivRM^*

Firstly, we introduce the output randomization technique: users whose values fall within the range, they normalize, perturb and then denormalize their values. And others draw a random value from the output domain of the NVP mechanism. However, when a perturbation mechanism produces unbounded noise (e.g., LM), it is not feasible to randomly select values from real number field $\mathbb{R}$. As such, we categorize them into two cases, namely bounded and unbounded LDP mechanisms, and discuss them as follows (Assuming that the privacy budget allocated for Phase 2 is ϵ).

Case 1: Bounded LDP Mechanisms. As for the bounded NVP mechanisms, the details of Phase 2 remain consistent with those mentioned earlier. Here, we take SR as an example to illustrate the process. When using SR, for users within the range, they directly perturb their values with the sensitivity $|r - l|$ through SR. For users outside the range, each of their outputs satisfies the Bernoulli distribution with probability 0.5 in $\mathbb{M}_{\text{SR}}(\epsilon) = \{-C_{\text{SR}}(\epsilon), C_{\text{SR}}(\epsilon)\}$.

Case 2: Unbounded LDP Mechanisms. As for unbounded LDP mechanisms, the idea is to truncate their output domains, which allows us to discuss the unbounded case by transforming it into the bounded case. Here, we take LM as an example. Without loss of generality, we truncate the output domain to the input domain, $[-1, 1]$. For users within the range, they perturb their values by LM and then truncate the perturbed data to $[-1, 1]$. As for out-of-range users, we take the input value of 0 as the reference standard. When the input is 0, the sum of probabilities at truncation points is $e^{-\epsilon}$, and the sum of probabilities within $(-1, 1)$ is $1 - e^{-\epsilon}$. Therefore, they should follow this reporting strategy: choosing a value uniformly in the range $(-1, 1)$ with probability $\alpha_{\text{TLM}}(\epsilon) = 1 - e^{-\epsilon}$, and randomly selecting a value from $\{-1, 1\}$ with probability $1 - \alpha_{\text{TLM}}(\epsilon) = e^{-\epsilon}$. This method is designated as Truncation-based Laplace Mechanism (TLM).

THEOREM 3. *For TLM, if the truncation range is $[-1, 1]$, it achieves ϵ -LDP.*

PROOF. The proof appears in the full version [3]. □

After introducing these two cases, we can apply output randomization to all NVP mechanisms. And it is clear that $\mathbb{E}[\hat{t}_{\text{out}}]$ remains constant, i.e., $\mathbb{E}[\hat{t}_{\text{out}}] = \frac{l+r}{2}$.

Note that PrivRM^l estimates both the range count and mean across all users, necessitating the division of the privacy budget ϵ into two parts. To avoid this division of the privacy budget, we could consider estimating the range count and mean using either in-range or out-of-range values exclusively. Specifically, given a range $[l, r]$, those in-range values could serve for estimating the range sum $s_{(l,r)}$, while out-of-range values could contribute to the range count $n_{(l,r)}$. For the later, it is feasible since the total number of in-range and out-of-range values always equals n . This idea directly leads to our optimized implementation, PrivRM^* , of the framework.

The details of PrivRM^* are shown in Algorithm 2. Firstly in Phase 1, the out-of-range users apply RR perturbation with probability p ($p > 0.5$) (Line 5), while others report 0 or 1 uniformly

Algorithm 2: Workflow of *PrivRM**

Input: A set of user values $\{t_1, t_2, \dots, t_n\}$, query range $[l, r]$, NVP mechanism $\mathcal{A}(\cdot)$, output domain $\mathbb{M}_{\mathcal{A}}(\cdot)$, perturbation probability p

Output: The estimated range mean $\hat{m}_{(l,r)}$

```

1  for  $i = 1$  to  $n$  do                                     // Phase 1
2      if  $t_i \in [l, r]$  then
3          | Report bit 1 or 0 uniformly at random
4      else                                                 // RR Perturbation
5          | Report bit 0 (resp. 1) with probability  $p$  (resp.  $1 - p$ )
6  Aggregator derives the count  $\hat{n}_{(l,r)} = \frac{2\#('1') + 2(p-1)n}{2p-1}$ 
7  Set  $\epsilon' = \log(\frac{p}{1-p})$ 
8  for  $i = 1$  to  $n$  do                                     // Phase 2
9      if  $t_i \in [l, r]$  then
10         | Normalization, Perturbation, Denormalization with  $\epsilon'$ 
11     else
12         | Randomly draw a value  $y_i \in \mathbb{M}_{\mathcal{A}}(\epsilon')$ 
13         | Denormalization
14 Calculate  $\hat{m}_{(l,r)}$  by Equation 6, where  $\mathbb{E}[\hat{t}_{out}] = \frac{l+r}{2}$ 
15 return  $\hat{m}_{(l,r)}$ 

```

at random (Line 3). Then the collector aggregates the range count $\hat{n}_{(l,r)}$ (Line 6). Note that the aggregation in Line 6 is different from that in *PrivRM^I* (Line 7 of Algorithm 1), since we differentiate the perturbation between in-range and out-of-range values. The following Theorem 4 establishes the correctness of this estimation.

In Phase 2, users' reporting way is output randomization, but the difference is that the privacy budget ϵ' used in this phase depends on p . On the other hand, since the NVP mechanism needs to take a privacy budget as input, we empirically set it as $\epsilon' = \log(\frac{p}{1-p})$ (Line 7). The privacy guarantee will be further analyzed in Section 5.2. For values within the range, they will go through normalization, perturbation and denormalization with privacy budget ϵ' (Lines 9-10). As for the values outside the range, they will be randomized into a value drawn from the output domain $\mathbb{M}_{\mathcal{A}}(\epsilon')$ (Line 12). Ultimately, the collector estimates the range mean $\hat{m}_{(l,r)}$ via Equation 6 (Line 14).

THEOREM 4. In Algorithm 2, we have $\mathbb{E}[\frac{2\#('1') + 2(p-1)n}{2p-1}] = n_{(l,r)}$.

PROOF. Bit 1 will be reported from out-of-range values with probability $1 - p$ or in-range values with probability 0.5. Thus, the expected number of observed bit 1 is

$$\mathbb{E}[\#('1')] = (1 - p)(n - n_{(l,r)}) + 0.5n_{(l,r)}. \quad (8)$$

Through Equation 8, we have,

$$\mathbb{E}[\frac{2\#('1') + 2(p-1)n}{2p-1}] = n_{(l,r)}.$$

□

5.2 Privacy Analysis of *PrivRM**

Given a perturbation probability p , the following Theorem 5 demonstrates that we can apply various NVP mechanisms to *PrivRM** under LDP setting.

Table 2. Relationship between p and ϵ for different NVP mechanisms

Mechanism $\mathcal{A}(\cdot)$	Privacy budget achieved by PrivRM^* with $\mathcal{A}(\cdot)$
TLM	$\epsilon = \begin{cases} \log\left(\frac{4p\left(\sqrt{\frac{p}{1-p}}p + p - \sqrt{\frac{p}{1-p}}\right)}{(1-p)\log\left(\frac{p}{1-p}\right)}\right), & \text{if } p < 0.75, \\ \log\left(\frac{1}{2-2p}\sqrt{\frac{p}{1-p}}\right), & \text{otherwise.} \end{cases}$
SR	$\epsilon = \log\left(\frac{p}{1-p}\right)$
PM	$\epsilon = \log\left(\frac{1}{2-2p}\sqrt{\frac{p}{1-p}}\right)$
HM	$\epsilon = \begin{cases} \log\left(\frac{p}{1-p}\right), & \text{if } p < 0.65, \\ \log\left(\frac{1}{2-2p}\sqrt{\frac{p}{1-p}}\right), & \text{otherwise.} \end{cases}$
Unbiased SW	$\epsilon = \log\left(\frac{2p-1}{2(p-1)^2\log\left(\frac{p}{1-p}\right)}\right)$

THEOREM 5. Given a probability p , PrivRM^* satisfies ϵ -LDP, where $\epsilon = \log(\max\left\{\frac{p}{1-p}, \frac{p \cdot q}{0.5 \cdot q'}, \frac{0.5 \cdot p'}{(1-p) \cdot q}\right\})$, p' and q' are the minimum and maximum in the perturbation probability density function (probability mass function for SR) of the NVP mechanism, and q is the probability of a random output on $\mathbb{M}_{\mathcal{A}}(\cdot)$.

PROOF. For an input value t_i , let $s_i = \langle \hat{k}_i, \hat{t}_i \rangle$ be the output of PrivRM^* , where $\hat{k}_i$ and $\hat{t}_i$ correspond to the outputs of Phase 1 and Phase 2, respectively. According to Definition 1, for any pair of distinct input values, the probability ratio for producing the same output must fall between $e^{-\epsilon}$ and e^{ϵ} . Here we consider two distinct cases: the first is when both users with differing inputs are within the query range, or alternatively, both are outside this range. The second case occurs when one user is inside the range and another is outside.

Case 1: If both values t_1, t_2 come from outside the query range, then in Phase 2, output randomization occurs for both. For any output s we have,

$$\frac{\Pr[s \mid t_1]}{\Pr[s \mid t_2]} = \frac{p \cdot q}{p \cdot q} = \frac{(1-p) \cdot q}{(1-p) \cdot q} = 1.$$

If both values t_1, t_2 are inside the query range $[l, r]$, they both report 0/1 randomly during Phase 1. Similarly they both use NVP perturbations with a privacy budget of $\epsilon' = \log\left(\frac{p}{1-p}\right)$ in Phase 2. For any output s we have,

$$\frac{\Pr[s \mid t_1]}{\Pr[s \mid t_2]} \leq \frac{0.5 \cdot p'}{0.5 \cdot q'} = e^{\epsilon'} = \frac{p}{1-p}.$$

Case 2: Assume that input $t_1 \in [l, r]$ and $t_2 \notin [l, r]$, it is easy to see that

$$\frac{0.5 \cdot q'}{p \cdot q} \leq \frac{\Pr[s \mid t_1]}{\Pr[s \mid t_2]} \leq \frac{0.5 \cdot p'}{(1-p) \cdot q}.$$

Therefore, according to Definition 1, PrivRM^* satisfies ϵ -LDP, where $\epsilon = \log(\max\left\{\frac{p}{1-p}, \frac{p \cdot q}{0.5 \cdot q'}, \frac{0.5 \cdot p'}{(1-p) \cdot q}\right\})$. $\square$

In what follows, we further analyze the privacy guarantee when integrating TLM, SR, PM, HM and Unbiased SW into PrivRM^* , as shown in Theorems 6-10. For ease of reference, in Table 2, we summarize the relationship between the perturbation probability p and the privacy budget ϵ achieved by PrivRM^* .

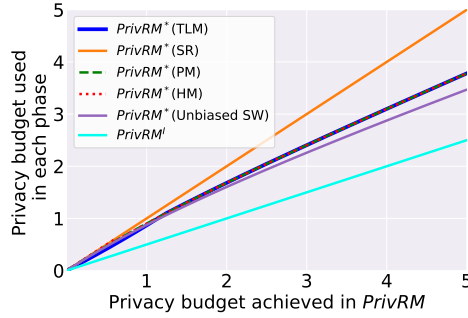


Fig. 2. Corresponding privacy parameter used in each phase - Privacy budget achieved in *PrivRM*

THEOREM 6. Applying TLM to *PrivRM*^{*} satisfies ϵ -LDP, where

$$\epsilon = \begin{cases} \log\left(\frac{4p\left(\sqrt{\frac{p}{1-p}}p + p - \sqrt{\frac{p}{1-p}}\right)}{(1-p)\log\left(\frac{p}{1-p}\right)}\right), & \text{if } p < 0.75, \\ \log\left(\frac{1}{2-2p}\sqrt{\frac{p}{1-p}}\right), & \text{otherwise.} \end{cases}$$

PROOF. The proof appears in the full version [3]. □

THEOREM 7. Applying SR to *PrivRM*^{*} satisfies ϵ -LDP, where $\epsilon = \log\left(\frac{p}{1-p}\right)$.

PROOF. The proof appears in the full version [3]. □

THEOREM 8. Applying PM to *PrivRM*^{*} satisfies ϵ -LDP, where $\epsilon = \log\left(\frac{1}{2-2p}\sqrt{\frac{p}{1-p}}\right)$.

PROOF. The proof appears in the full version [3]. □

THEOREM 9. Applying HM to *PrivRM*^{*} satisfies ϵ -LDP, where

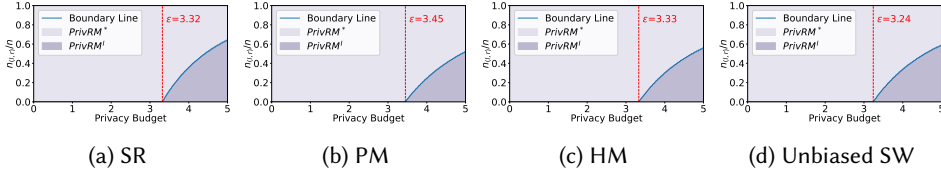
$$\epsilon = \begin{cases} \log\left(\frac{p}{1-p}\right), & \text{if } p < 0.65, \\ \log\left(\frac{1}{2-2p}\sqrt{\frac{p}{1-p}}\right), & \text{otherwise.} \end{cases}$$

PROOF. The proof appears in the full version [3]. □

THEOREM 10. Applying Unbiased SW to *PrivRM*^{*} satisfies ϵ -LDP, where $\epsilon = \log\left(\frac{2p-1}{2(p-1)^2\log\left(\frac{p}{1-p}\right)}\right)$.

PROOF. The proof appears in the full version [3]. □

Now given a privacy budget ϵ , we can compute the privacy parameters used in the two phases of *PrivRM*^{*}, including perturbation probability p and the privacy budget ϵ' used in Phase 2. With varying privacy budgets eventually achieved by *PrivRM*^{*}, Figure 2 shows the corresponding privacy budget actually used in each phase, when integrating each NVP mechanism. For comparison purpose, we also plot the cyan curve, the privacy budget used in each phase of *PrivRM*^l (i.e., $\epsilon/2$). Obviously, we can observe that regardless of the NVP mechanism used, the privacy budget used by *PrivRM*^{*} is always larger than that used by another implementation. This means that, to achieve a same privacy level, *PrivRM*^{*} can use a larger privacy budget, leading to more accurate results. we also observe that when $\epsilon > 1.24$, the privacy parameters in TLM, PM and HM are the same. The reason is that these three mechanisms satisfies $\log\left(\frac{1}{2-2p}\sqrt{\frac{p}{1-p}}\right)$ -LDP in this case.

Fig. 3. Comparison between $PrivRM^*$ and $PrivRM^I$

5.3 Which Implementation to Use?

Note that all the five NVP mechanisms can be integrated into our implementations, namely $PrivRM^I$ and $PrivRM^*$. We may need to answer a question: which implementation to use in a given setting? Fortunately, Theorem 2 enables us to theoretically evaluate these three implementations through estimation variance.

First we need to calculate the variance of $PrivRM^*$. With the above Theorems 6-10 and Theorem 2, given the privacy budget ϵ , we can compute the corresponding the variance of $PrivRM^*$. Let $X = \sum_{i=1}^n \hat{t}_i - n\mathbb{E}[\hat{t}_{out}]$ and $Y = \hat{n}_{(l,r)}$. When the NVP mechanism is TLM, truncation similarly results in the equality of $\mathbb{E}[X]$ and $s_{(l,r)} + (n - n_{(l,r)})\frac{l+r}{2}$ only if the data within the query range follows a symmetric distribution. When employing other NVP mechanisms, $PrivRM^*$ is also unbiased in both phases. Thus,

$$\mathbb{E}[X] = s_{(l,r)} - n_{(l,r)} \left(\frac{l+r}{2} \right), \quad \mathbb{E}[Y] = n_{(l,r)}.$$

In $PrivRM^*$, in-range and out-of-range users report randomly in Phase 1 and Phase 2 respectively, so we have

$$Var(X) = n_{(l,r)} Var(\epsilon') + (n - n_{(l,r)})U(\epsilon'),$$

$$Var(Y) = \frac{n_{(l,r)}p(1-p) + 0.25(n - n_{(l,r)})}{(2p-1)^2}.$$

Subsequently, the variance of $PrivRM^*$ can be calculated through Equation 7.

After obtaining the variance of $PrivRM^*$, we can choose the optimal implementation. When the data within the range is not symmetrically distributed, TLM is biased in $PrivRM^*$. Therefore, under these conditions, $PrivRM^I$ should be used. When other mechanisms are employed, the sum and count estimation of these two implementations are unbiased, that is, $\mathbb{E}[X]$ and $\mathbb{E}[Y]$ are the same. Besides, since n is much larger than $Var(Y)$, we have $\frac{Var(Y)}{\mathbb{E}^2[Y]} \approx 0$. This means that the final variance depends on $Var(X)$. Now, we compare $PrivRM^I$ and $PrivRM^*$. As we mentioned above, given the privacy budget ϵ , we can obtain the corresponding ϵ' . The variance of $PrivRM^*$ is less than that of $PrivRM^I$ when the following conditions are satisfied:

$$\Delta = nVar(\epsilon) - n_{(l,r)}Var(\epsilon') - (n - n_{(l,r)})U(\epsilon') > 0.$$

Black-box Guidelines. Given privacy budget ϵ , $n_{(l,r)}$ and NVP mechanism we used $\mathcal{A}(\cdot)$ as inputs, our analysis gives the following guidelines for choosing optimal implementation.

- When employing the LM, it is advisable to use $PrivRM^I$, as it remains unbiased under all circumstances.
- When using other mechanisms, if $\Delta > 0$, we should employ the $PrivRM^*$; otherwise, $PrivRM^I$ should be used.

Discussion. We can change inputs of black-box guidelines to get the implementation with the lowest variance¹, as shown in Figure 3. It is evident that when the privacy budget is relatively small ($\epsilon < 3.32$ for SR, $\epsilon < 3.45$ for PM, $\epsilon < 3.33$ for HM, $\epsilon < 3.24$ for Unbiased SW), the variance of *PrivRM** is consistently minimized. Intuitively, when the privacy budget is small, the variance of any mechanism decreases exponentially as the privacy budget increases. This is why the *PrivRM** performs better when the privacy budget is small.

6 Distribution-aware Optimization on Framework

Given a specific range, the *PrivRM* framework reduces the sensitivity of estimating the range mean from the entire domain to the query range. Upon looking into the range, it's observed that for a skewed data distribution, where the majority of values are concentrated in an area much smaller than the query range, the noise sensitivity can be further reduced by confining the perturbation space to a densely populated region. Inspired by this, we propose a strategy of Adaptive Adjustment (AA) on perturbation space based on a binary search.

6.1 Adaptive Adjustment Strategy

When the majority of values are concentrated in a small range, confining the perturbation space to this dense region can reduce the estimation variance, enhancing the utility of the estimation result. Considering the lack of prior information on the data distribution, a feasible solution is to allocate a portion of users to estimate this distribution, based on which we can traverse the query range to explore the optimal densely populated region. However, since this strategy requires traversing all sub-intervals within the query range, it suffers from low efficiency. To address this issue, we can then use a binary search to identify the densely populated region. Based on this intuition, we design a strategy of Adaptive Adjustment (AA) on perturbation space which can be integrated into *PrivRM* framework to enhance the accuracy of range mean estimation. In the AA strategy, we divide the perturbation space equally until we identify the most suitable dense region.

6.2 Iteration Condition

The essence of the AA strategy lies in the iteration condition, which can be constructed by jointly considering the variance and the data loss from confining the perturbation space, using the total squared error as the objective.

Given the query range $[l, r]$, privacy budget ϵ and the number of iterations k , let Ω_k be the current dense region with a width of $\frac{r-l}{2^k}$. In our framework *PrivRM*, the total variance for values perturbation is equal to $Var(\sum_{i=1}^n \hat{t}_i) = Var(X_k) = Var(\sum_{i=1}^n \hat{t}_i - n\mathbb{E}[\hat{t}_{out}])$. So the final variance is $Var(X) \left(\frac{r-l}{2^k}\right)^2$. Beside, the bias error of data loss caused by confining the perturbation space is $(\sum_{t_i \notin \Omega_k \cap t_i \in [l, r]} t_i)^2$. So the total squared error is as follows:

$$\mathbb{E}_k^1 = Var(X_k) \left(\frac{r-l}{2^k}\right)^2 + \left(\sum_{t_i \notin \Omega_k \cap t_i \in [l, r]} t_i\right)^2. \quad (9)$$

Here we assume that the dense region after the k -th round of binary division is Ω_{k+1} , from which we can easily infer that its width is $\frac{r-l}{2^k}$. Thus, we have,

¹Here to maximize the difference between the variance of NVP mechanism and that of uniform distribution, we use the best-case variance of the NVP mechanism.

Algorithm 3: Workflow of AA

Input: A set of user values $\{t_1, t_2, \dots, t_n\}$, query range $[l, r]$, NVP mechanism $\mathcal{A}(\cdot)$, privacy budget ϵ
Output: The dense region Ω

- 1 Randomly sample 10% of the values from $\{t_1, t_2, \dots, t_n\}$
- 2 Estimate the distribution of the sampled data by SW mechanism
- 3 Initialize $\Omega_1 = [l, r]$
- 4 **for** $k = 1$ **to** 10 **do** // $2^{10} = 1024$
- 5 Divide Ω_k equally into two sub-ranges Ω_k^1 and Ω_k^2
- 6 **if** $\sum \hat{d}_{\Omega_k^1}^i > \sum \hat{d}_{\Omega_k^2}^i$ **then**
- 7 | $\Omega_{k+1} = \Omega_k^1$
- 8 **else**
- 9 | $\Omega_{k+1} = \Omega_k^2$
- 10 Calculate the total squared error gap δ_k by Equation 11
- 11 **if** $\delta > 0$ **then** // Iteration condition
- 12 | $\Omega_k = \Omega_{k+1}$
- 13 **else**
- 14 | **return** Ω_k

$$\mathbb{E}_k^2 = \text{Var}(X_{k+1}) \left(\frac{r-l}{2^{k+1}} \right)^2 + \left(\sum_{t_i \notin \Omega_{k+1} \cap t_i \in [l, r]} t_i \right)^2. \quad (10)$$

Through Equation 9 and Equation 10, we can derive the reduction of total squared error as,

$$\delta_k = \mathbb{E}_k^1 - \mathbb{E}_k^2, \quad (11)$$

When $\delta_k > 0$, it indicates that the total squared error has decreased after the binary division, thereby justifying the progression to the $(k + 1)$ -th round of binary search.

Intuitively, in the k -th round of searching, if most users are concentrated in region Ω_{k+1} , we have $\text{Var}(X_k) \approx \text{Var}(X_{k+1})$ and $(\sum_{t_i \notin \Omega_k \cap t_i \in [l, r]} t_i)^2 \approx (\sum_{t_i \notin \Omega_{k+1} \cap t_i \in [l, r]} t_i)^2$. Under these circumstances, it becomes clear that $\mathbb{E}_k^1 \approx 4\mathbb{E}_k^2$, i.e., $\delta_k > 0$ always holds. This means that we can designate Ω_{k+1} as the current dense region and proceed to the next round of the search.

6.3 Putting Things Together

By putting things together, Algorithm 3 shows the workflow of AA strategy. We start by randomly sampling 10% of values for estimating the data distribution (Lines 1 and 2) and then divide the perturbation space. Here we take the first round as an example. In the first round, we first divide perturbation space $[l, r]$ equally into two sub-ranges $\Omega_1^1 = [l, (l+r)/2]$ and $\Omega_1^2 = [(l+r)/2, r]$ (Line 5). We then calculate the sum of the probability densities of these two sub-ranges separately, i.e., $\sum \hat{d}_{\Omega_1^1}^i$ and $\sum \hat{d}_{\Omega_1^2}^i$. We then determine which sub-range has a higher density and label it as the dense region (Lines 6-9). Subsequently, we calculate the total squared error gap δ_k (Line 10), and then assess whether the iteration condition holds (Line 11). If it proves to be favorable, binary search is further applied to the dense sub-range (Line 12). Otherwise, we return the current perturbation space and designate it as the final dense region (Line 16).

Beside, with the prior distribution, we can estimate the range count $n_{(l, r)}$. In turn, we can choose the optimal implementation through the black box in Section 5.4 based on our chosen NVP

Algorithm 4: Workflow of Multidimensional PrivRM^l

Input: A set of user values $\{T_1, T_2, \dots, T_n\}$, query range $[l_0, r_0]$, conditions $[l_k, r_k] (1 \leq k \leq d)$, NVP mechanism $\mathcal{A}(\cdot)$, privacy budget ϵ

Output: The estimated multidimensional range mean $\hat{m}_{(l_0, r_0)}$

```

1 for  $i = 1$  to  $n$  do // Phase 1
2   Set  $p = \frac{e^{\epsilon/2}}{1+e^{\epsilon/2}}$ 
3   if  $\forall k \in \{0, 1, \dots, d\}, t_i^k \in [l_k, r_k]$  then
4     Report bit 1 (resp. 0) with probability  $p$  (resp.  $1 - p$ )
5   else
6     Report bit 0 (resp. 1) with probability  $p$  (resp.  $1 - p$ )
7 Collector derives the range count  $\hat{n}_{(l_0, r_0)} = \frac{n(p-1) + \#('1')}{2p-1}$ 
8 for  $i = 1$  to  $n$  do // Phase 2
9   if  $\forall k \in \{0, 1, \dots, d\}, t_i^k \in [l_k, r_k]$  then
10    Normalization:  $t'_i = \frac{2(t_i^0 - l_0)}{r_0 - l_0} - 1$ 
11    Perturbation:  $y_i = \mathcal{A}(t'_i, \epsilon/2)$ 
12    Denormalization:  $\hat{t}_i^0 = \frac{(y_i+1)(r_0-l_0)}{2} + l_0$ 
13  else
14    Randomly draw a value  $t'_i \in [-1, 1]$ 
15    Perturbation:  $y_i = \mathcal{A}(t'_i, \epsilon/2)$ 
16    Denormalization:  $\hat{t}_i^0 = \frac{(y_i+1)(r_0-l_0)}{2} + l$ 
17 Calculate  $\hat{m}_{(l_0, r_0)}$  by Equation 6, where  $\mathbb{E}[\hat{l}_{out}] = \frac{l_0+r_0}{2}$ 
18 return  $\hat{m}_{(l_0, r_0)}$ 

```

mechanism and the privacy budget. Eventually, we can utilize $\text{Var}(X)$ of our chosen implementation to confine the perturbation space.

7 EXTENSIONS AND DISCUSSIONS

Following the introduction of our framework and strategy, in this section we will delve into the extension of regular range mean estimation to multidimensional cases. Besides, we will discuss the relationship between range mean estimation and general range query. Additionally, we also extend our Framework to multiple range mean estimations.

Multidimensional Range Mean Estimation. When the query involves multiple attributes, our methods can be naturally extended to handle such multidimensional range mean estimation. For instance, considering a target attribute C_0 and other d condition attributes $\{C_1, C_2, \dots, C_d\}$, user u_i possesses $d + 1$ numerical values, denoted as $T_i = \{t_i^0, t_i^1, \dots, t_i^d\}$. We aim to estimate the mean of values corresponding to attribute C_0 within the range $[l_0, r_0]$, while satisfying conditions that the other d attributes' values fall within ranges $[l_k, r_k] (1 \leq k \leq d)$ respectively. Algorithms 4 and 5 show the pseudo-code on how PrivRM^l and PrivRM^* can handle multidimensional queries respectively. Upon comparison, the main difference between single-dimensional PrivRM and multidimensional PrivRM lies in the judgment conditions: In the single-dimensional PrivRM , each user only considers whether his value falls within the given query range; however, in multidimensional PrivRM , each user determines whether his record meets the query conditions across all dimensions, and only processes value corresponding to target attribute C_0 in Phase 2.

Relationship with General Range Query. Answering range query under Local Differential Privacy is a widely studied problem [12, 24, 32]. The main difference between range mean estimation

Algorithm 5: Workflow of Multidimensional *PrivRM**

Input: A set of user values $\{T_1, T_2, \dots, T_n\}$, query range $[l_0, r_0]$, conditions $[l_k, r_k] (1 \leq k \leq d)$, NVP mechanism $\mathcal{A}(\cdot)$, output domain $\mathbb{M}_{\mathcal{A}}(\cdot)$, perturbation probability p

Output: The estimated multidimensional range mean $\hat{m}_{(l_0, r_0)}$

```

1 for  $i = 1$  to  $n$  do // Phase 1
2   if  $\forall k \in \{0, 1, \dots, d\}, t_i^k \in [l_k, r_k]$  then
3     | Report bit 1 or 0 uniformly at random
4   else // RR Perturbation
5     | Report bit 0 (resp. 1) with probability  $p$  (resp.  $1 - p$ )
6   Aggregator derives the count  $\hat{n}_{(l_0, r_0)} = \frac{2^{\#('1')} + 2(p-1)n}{2p-1}$ 
7   Set  $\epsilon' = \log(\frac{p}{1-p})$ 
8   for  $i = 1$  to  $n$  do // Phase 2
9     if  $\forall k \in \{0, 1, \dots, d\}, t_i^k \in [l_k, r_k]$  then
10      | Process  $t_i^0$  through Normalization, Perturbation, Denormalization with  $\epsilon'$ 
11    else
12      | Randomly draw a value  $y_i \in \mathbb{M}_{\mathcal{A}}(\epsilon')$ 
13      | Denormalization
14   Calculate  $\hat{m}_{(l_0, r_0)}$  by Equation 6, where  $\mathbb{E}[\hat{t}_{out}] = \frac{l_0 + r_0}{2}$ 
15 return  $\hat{m}_{(l_0, r_0)}$ 

```

and range query is the query type. Taking the single dimension as an example, given a query range (l, r) , the task of range query is to estimate the frequency of values within this range; however, the goal of range mean estimation is to estimate the mean of values within this range.

Multiple Range Mean Estimations. When data collector wants to query the mean values for multiple ranges, as exemplified by a given set of ranges $\{(l_i, r_i)\} (i = 1, 2, \dots, N)$, the objective is to estimate the mean value for each individual range $m_{(l_i, r_i)} (i = 1, 2, \dots, N)$. Based on Theorem 1, we can allocate the privacy budget evenly across each query, and then we can extend our framework to multiple range queries. Specially, given N range mean queries, we split the privacy budget ϵ among these queries equally, i.e., each with $\frac{\epsilon}{N}$. Then according to Theorem 2, we can theoretically derive the variance of the estimated range mean. As for the two baselines, the total budget ϵ will be used for perturbation. Since the variance heavily relies on the dataset distribution or the post-processing step, it is not feasible to directly derive their variance theoretically. Nevertheless, we can calculate their empirical variance. Then we can analyze the number of queries N our approaches could offer better utility than two baselines. Given the query range (l, r) , assuming that the smallest empirical variance of baselines under the current settings is Var_{base} , we can calculate N as follows:

$$\max N, \text{ subject to: } Var(\hat{m}_{(l, r)}) \leq Var_{base},$$

where $Var(\hat{m}_{(l, r)})$ is the theoretical variance of *PrivRM*.

8 Experimental Evaluation

In this section, we evaluate the performance of our proposed framework *PrivRM* for range mean estimation.

8.1 Experiment Setup

Datasets. We conduct experiment over three real-world datasets. Table 3 gives the distribution statistics of the datasets. Note that the original value domains of them are different, and we normalize all the values into $[-1, 1]$.

Table 3. Dataset Distribution Characteristics

Datasets	n	Avg.	Med.	Dense region (99%)	Normalized dense region
Kosarak	990002	8.1	3.0	(1, 89)	(-1.0, -0.9957)
House	2226382	523832.7	325000	(0, 3899000)	(-1.0, -0.9963)
Fare	7696617	12.5	9.0	(2.5, 52)	(-0.9988, -0.9986)

- Kosarak [2] is click-stream data, where the categories of clicks are considered representative values for the users. The value ranges from 1 to 41, 270, with a total of 990, 002 samples.
- House [4] contains real estate listings in the US. The value ranges from 0 to approximately 2.15×10^9 , with a total of 2, 226, 382 samples.
- Fare [1] contains fare data for taxi travel in New York City, specifically the Yellow Taxi Trip Records from January 2019. The value ranges from -362 to 623, 259, with a total of 7, 696, 617 samples.

Competitors. In our experiments, we first implement two baseline methods introduced in Section 3.2, namely Direct Estimation and Distribution-based Estimation. As for our methods, we focus on two versions. The first is *Optimal-PrivRM*, the best choice from *PrivRM^l* and *PrivRM^{*}*, according to the black-box guideline described in Section 5.3. The second is *Optimal-PrivRM-AA*, the *Optimal-PrivRM* coupled with AA strategy.

Experiment Design. We evaluate the performance of the four competitors with a fixed query range or with a fixed privacy budget, evaluate Distribution-based Estimation with different bin sizes, test the effectiveness and costs of AA strategy, show the results of multiple range mean queries, verify the correctness of our black-box selection scheme, and compare the performance of Original SW and Unbiased SW for mean estimation.

Metrics. As for the result evaluation, we employ the metric of Mean Squared Error (MSE)[27], which quantifies the average squared difference between the estimated values $\hat{m}_{(l,r)}$ and the actual values $m_{(l,r)}$. Formally,

$$MSE(l, r) = \frac{1}{N} \sum_N [\hat{m}_{(l,r)} - m_{(l,r)}]^2, \quad (12)$$

where N represents the number of repetitions for each experiment. In our study, N is set to 50.

We conduct experiments using Python 3.11.5 and the Numpy 1.24.3 library on a desktop equipped with an Intel Core i5-13400F 1.50 GHz CPU and 64GB of RAM, running Windows 11.

8.2 Impact of Privacy Budgets

In this section, we evaluate the performance of four methods across four datasets and different NVP mechanisms with varying privacy budgets. As for the query range $[l, r]$, we set it to the first half, i.e. $[-1, 0]$. It is worth noting that only those NVP mechanisms whose output domains are continuous can be applied to Direct Estimation. Therefore, Direct Estimation in SR and HM is not feasible in Figure 4.

Overall, we observe that as the privacy budget increases, the MSE of all methods decreases. Besides, *Optimal-PrivRM-AA* performs the best, followed by *Optimal-PrivRM*, and finally two baseline methods. It is obvious that the gap between *Optimal-PrivRM-AA* and *Optimal-PrivRM* widens as the privacy budget increases. This is because, when the query range is fixed, the sensitivity of *Optimal-PrivRM* remains constant. As the privacy budget increases, the AA strategy's prior distribution estimate becomes more accurate, reducing the adjusted perturbation space and resulting in less noise in *Optimal-PrivRM-AA*. Additionally, the performance of Direct Estimation does not significantly differ when using LM, PM, and Unbiased SW, regardless of the privacy budget. This

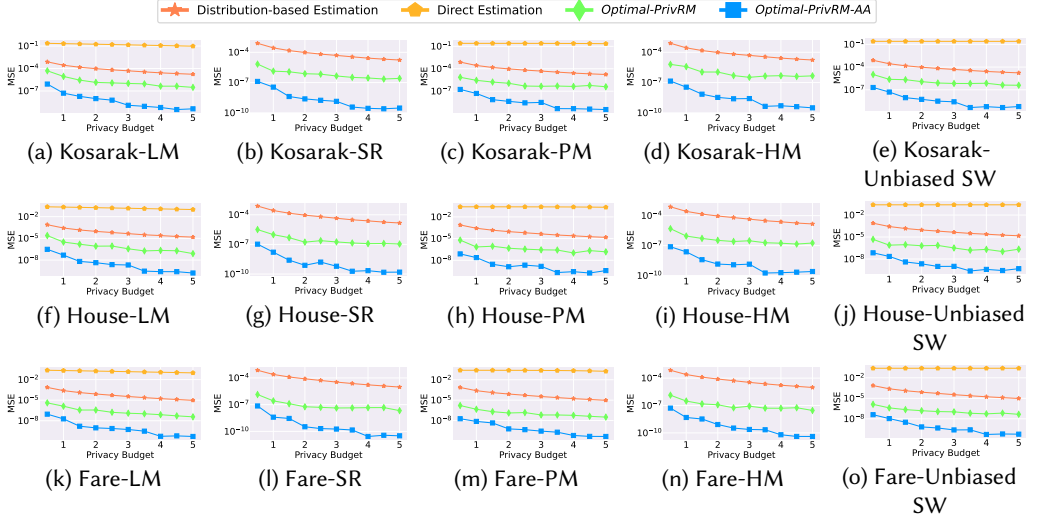


Fig. 4. Performance of different methods with varying privacy budgets.

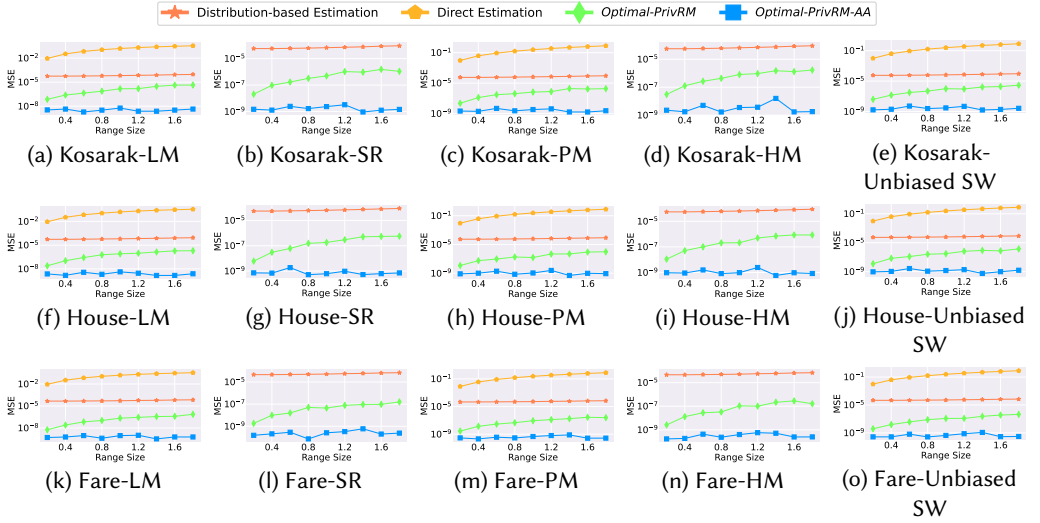


Fig. 5. Performance on real-world dataset with varying size of query range. Query range starts from -1.

is because, for a given query range, the query bias is fixed. As the privacy budget increases, the variance of the estimated result decreases, but this change is negligible compared to the bias.

8.3 Impact of Range Size

This subsection studies the impact of the range size. The results are shown in Figure 5, where the range size refers to the span length starting from -1 , and the privacy budget is fixed at $\epsilon = 2.5$. Likewise, results of Direct Estimation in SR and HM are not feasible.

We observe that *Optimal-PrivRM-AA* always yields the highest accuracy, followed by *Optimal-PrivRM*. Notably, with an increasing range size, the MSE of *Optimal-PrivRM* also increases, whereas MSE of *Optimal-PrivRM-AA* remains almost unchanged. This can be attributed to the lower noise

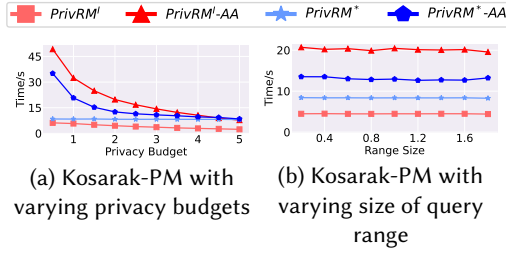


Fig. 6. Computation time for implementations with and without AA strategy.

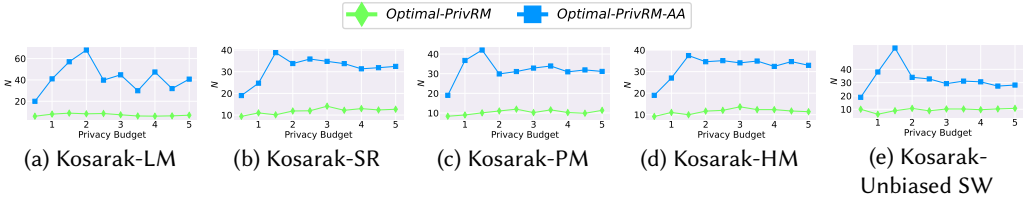


Fig. 7. Results of multiple range mean queries on Kosarak Dataset with varying privacy budgets. The query range is $[-1, 0]$.

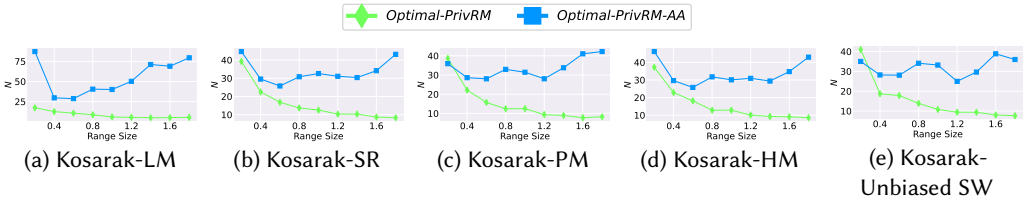


Fig. 8. Results of multiple range mean queries on Kosarak Dataset with varying size of query range. The privacy budget is 2.5. Query range starts from -1

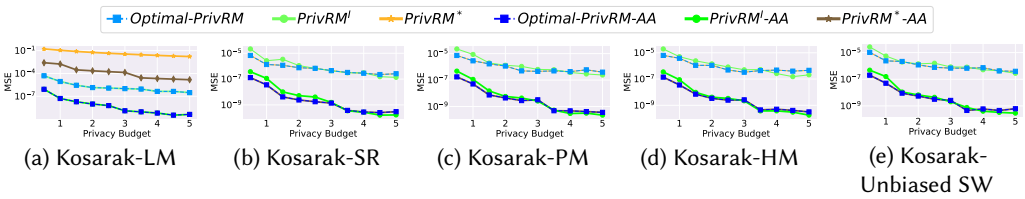


Fig. 9. Results of different implementations on Kosarak Dataset. The query range is $[-1, 0]$.

sensitivity of *Optimal-PrivRM-AA* than *Optimal-PrivRM*. The sensitivity of *Optimal-PrivRM* increases with the query range size, whereas *Optimal-PrivRM-AA* always tunes the query range to the densely distributed region, resulting in a relatively constant sensitivity.

8.4 Effectiveness and Costs of AA Strategy

From the above two subsections, it is evident that the AA strategy indeed enhances the accuracy of queries. In this subsection, we directly validate the effectiveness of AA strategy through adjustments to the query range facilitated by it. Additionally, we will also assess the extra costs associated with the AA strategy (i.e., time consumption cost).

Table 4. The theoretical N_{LM} on Kosarak with varying privacy budgets.

Privacy Budget	$\epsilon = 1.0$	$\epsilon = 2.0$	$\epsilon = 3.0$	$\epsilon = 4.0$	$\epsilon = 5.0$
N_{LM}	6	7	7	7	7

Table 5. The theoretical N_{LM} on Kosarak with varying range size. Query range starts from -1.

Range Size	$ r - l = 0.2$	$ r - l = 0.6$	$ r - l = 1.0$	$ r - l = 1.4$	$ r - l = 1.8$
N_{LM}	32	11	7	6	5

We design experiments on the Kosarak dataset, with Tables 6 and 7 respectively presenting the adjusted query ranges under different privacy budgets and size of query range as estimated by the AA strategy. From Table 1, we observe that although the query range is initially set to $[-1, 0]$, the range adjusted by the AA strategy noticeably shifts closer to the actual dense region provided in Table 3. Additionally, It is evident that as the privacy budget increases, the adjusted range aligns more closely with the actual dense region. This phenomenon can be attributed to the fact that a larger privacy budget allows for a more accurate estimation of prior knowledge. Furthermore, Table 7 reveals that as the query range varies, the dense distribution ranges estimated by the AA strategy remain relatively close to the actual dense region.

Regarding the additional costs associated with the AA strategy, Figure 6 illustrates the computational time costs for implementations with and without it under various privacy budgets and query range sizes. Through comparison, we observe that implementations incorporating the AA strategy exhibit longer computational times than those without it. The disparity in computational time is more pronounced with smaller privacy budgets and appears to be minimally influenced by the size of query range. This is due to the fact that implementations employing the AA strategy necessitate the application of SW for estimating the prior distribution. During this estimation, the smaller the privacy budget, the longer time required for Expectation Maximization post-processing. Overall, irrespective of whether the implementation incorporates AA strategy or not, our framework achieves range mean estimation with little time expenditure.

8.5 Results of Multiple Range Mean Queries

We conduct experiments on Kosarak dataset. Here we take LM as an example and present corresponding theoretical N in Tables 4 and 5. Additionally, we evaluate the results of multiple range mean queries under various NVP mechanisms, as shown in Figures 7 and 8. We can observe that the theoretically optimal N and the empirical N are very close. Overall, in most cases, the number of queries N for *Optimal-PrivRM* is larger than 10. When the AA strategy is incorporated, the number of queries also depends on the dense regions estimated by the AA strategy. Although providing corresponding theoretical analysis is infeasible, experiments reveal that the number of queries N for *Optimal-PrivRM-AA* is consistently larger than 20.

As shown in Figures 7 and 8, we also observe that the number of queries in *Optimal-PrivRM-AA* does not exhibit a monotonic trend. As the privacy budget increases, the number of queries initially rises and then stabilizes at a certain level. This phenomenon can be explained as follows: when the privacy budget is relatively small, the AA strategy is prone to data loss due to inaccuracies in the prior distribution, leading to an increase in the number of queries as the budget grows. As the privacy budget continues to increase, the data loss caused by the AA strategy gradually decreases. At this stage, the impact of privacy budget growth on both the baseline and *Optimal-PrivRM-AA* becomes similar, resulting in the number of queries stabilizing at approximately 30. As the query range size increases, the number of queries in *Optimal-PrivRM-AA* initially decreases and then rises.

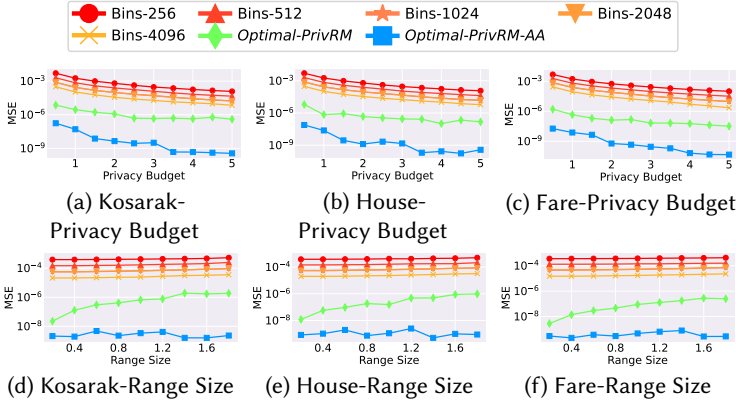


Fig. 10. Performance with varying bins of distribution. The NVP mechanism used is PM. The query range and privacy budget are the same as those set in Figures 4 and 5.

This pattern can be attributed to the following factors: when both the query range and privacy budget are small, the AA strategy tends to fail, making the performance of *Optimal-PrivRM-AA* and Method *Optimal-PrivRM* quite similar, as illustrated in Figures 8(b)-(e). As the query range expands, the data loss caused by the AA strategy diminishes, consequently leading to an increase in the number of queries. Additional experimental details are provided in the full version [3].

8.6 Correctness of Black-box Guidelines

In this subsection, we validate the correctness of the black-box guidelines proposed in Section 5.3. Given a privacy budget and an NVP mechanism, it guides us to choose the optimal implementation of the *PrivRM* framework from *PrivRM^l* and *PrivRM^{*}*.

We conduct experiment on Kosarak dataset, and show the results in Figure 9. First, we analyze the implementations without AA strategy, namely *PrivRM^l*, *PrivRM^{*}* and *Optimal-PrivRM*. We observe that *Optimal-PrivRM* achieves the lowest MSE in most cases.

In Figure 9(a), the lines for *Optimal-PrivRM* and *PrivRM^l* overlap. This is because when the NVP mechanism is LM, *PrivRM^l* is optimal according to the black box. When applying other mechanisms, black box indicates that *PrivRM^{*}* is optimal, so the line for *Optimal-PrivRM* overlaps with that of *PrivRM^{*}*. However, when the privacy budget exceeds 3.5, we observe that *Optimal-PrivRM* may not always be the best. This is because when $\epsilon > 3.5$, the optimal implementation starts to depend on the range count according to the prior distribution estimation. However, the prior distribution estimation involves noise, leading to inaccurate range counts. For a detailed theoretical analysis, please refer to Section 5.3.

Integrating our black-box guideline into *PrivRM-AA* leads to a solution *Optimal-PrivRM-AA*, which achieves the best performance in most cases. *Optimal-PrivRM-AA* overlaps with *PrivRM^l-AA* in Figure 9(a), and overlaps with *PrivRM^{*}-AA* in Figures 9(b)-(d), which is consistent with the non-AA versions. This is because, for a given dataset and query range, the range counts of the adjusted perturbation space derived from the AA strategy are highly consistent with those of the original ranges.

8.7 Impact of the Number of Bins on Distribution-based Estimation

For baseline method Distribution-based Estimation, more bins ensures more accurate results, but at the cost of more time consumption. To show some empirical analysis, we set the number of bins to

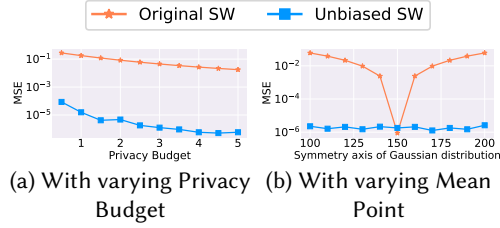


Fig. 11. Comparison of Original SW and Unbiased SW

Table 6. Most frequent dense regions estimated by AA strategy with varying privacy budgets on Kosarak. Query range is $[-1, 0]$.

Privacy Budget	$\epsilon = 1.0$	$\epsilon = 2.0$	$\epsilon = 3.0$	$\epsilon = 4.0$	$\epsilon = 5.0$
$\text{PrivRM}^I\text{-AA}$	$(-1, -0.875)$	$(-1, -0.9375)$	$(-1, -0.9375)$	$(-1, -0.96875)$	$(-1, -0.96875)$
$\text{PrivRM}^e\text{-AA}$	$(-1, -0.875)$	$(-1, -0.9375)$	$(-1, -0.9375)$	$(-1, -0.96875)$	$(-1, -0.96875)$

Table 7. Most frequent dense regions estimated by AA strategy with varying size of query range on Kosarak. Privacy budget is 2.5. Query range starts from -1.

Range Size	$ r - l = 0.2$	$ r - l = 0.6$	$ r - l = 1.0$	$ r - l = 1.4$	$ r - l = 1.8$
$\text{PrivRM}^I\text{-AA}$	$(-1, -0.95)$	$(-1, -0.9625)$	$(-1, -0.9375)$	$(-1, -0.95625)$	$(-1, -0.94375)$
$\text{PrivRM}^e\text{-AA}$	$(-1, -0.95)$	$(-1, -0.925)$	$(-1, -0.9375)$	$(-1, -0.95625)$	$(-1, -0.94375)$

256, 512, 1024, 2048, and 4096 in our evaluations, and utilize PM as the NVP mechanism for *Optimal-PrivRM* and *Optimal-PrivRM-AA*. As shown in Figure 10, the MSE of Distribution-based Estimation decreases as the number of bins increases. This is because more bins allows SW to estimate the changes in data distribution more accurately, which in turn leads to a performance improvement. Nevertheless, even when the number of bins is as large as 4096², Distribution-based Estimation still causes significantly higher MSE than that of *Optimal-PrivRM* and *Optimal-PrivRM-AA*.

8.8 Comparison between Original SW and Unbiased SW

As previously stated, the original SW mechanism is unbiased only if the data is symmetrically distributed over its value domain. Therefore, when the data is symmetrical, the mean estimation performance of the Original SW and the Unbiased SW should be similar. However, when the data is asymmetrical, the Original SW is expected to perform significantly worse than the Unbiased SW. To verify this hypothesis, we conduct experiments using both the Original SW and the Unbiased SW on Gaussian distributions with varying parameters.

In Figure 11(a), we use a Gaussian distribution centered at 100 with a domain of (0, 300), discarding any data point outside this range. This makes the data distribution asymmetric. We conduct experiments using both the Original SW and the Unbiased SW with varying privacy budgets. As demonstrated, the Original SW is less effective than the Unbiased SW for any privacy budget, as the Original SW introduces bias when estimating the mean.

Next, we fix the privacy budget and change the degree of asymmetry of the Gaussian distribution in the range by moving the symmetry axis from 100 to 200. As illustrated in Figure 11(b), the Original SW outperforms the Unbiased SW only when symmetry axis is 150, where the Gaussian distribution is symmetric over the domain (0, 300). Additionally, in this case, the variance of the Original SW is lower than that of the Unbiased SW, as shown in full version [3].

²In existing work [26], the number of bins is set to 1024.

9 Related Work

Differential privacy was initially introduced within a centralized framework [19]. To eliminate the need for a trusted data collector, the concept of local differential privacy (LDP) emerged, allowing individuals to independently apply perturbations to their data [16]. Over time, a variety of LDP solutions have been developed to address different statistical data collection challenges, including frequency estimation over categorical data [7, 23, 31] and mean estimation over numerical data [11, 30]. More recently, research in LDP has expanded to tackle more intricate tasks. These advanced applications include key-value data analysis [21, 40], heavy hitters identification [6, 8], trajectory data analysis [14, 41], marginal release [10, 42], graph data mining [28, 37], time series data collection [5, 38, 39] and range query [12, 22, 24, 32, 35]. Additionally, within the LDP framework, there are existing studies that separately provide strategies for attacking and defending these systems [13, 25, 34]. Moving forward, we will provide an overview of existing LDP research pertinent to our work, with a specific emphasis on mechanisms for handling categorical and numerical data.

LDP Mechanisms for categorical data A variety of LDP mechanisms have been designed specifically for perturbing categorical data. For binary data, the Randomized Response (RR) [33] serves as the simplest mechanism, while its extension, the Generalized Randomized Response (GRR) [23], addresses categories with a domain size greater than 2. To reduce the increasing perturbation noise associated with larger domain sizes in GRR, Wang et al. introduced the Optimized Unary Encoding (OUE) [31], which offers improved utility. Additionally, some other perturbation protocols such as RAPPOR [20], SHist [7], and subset selection [36] have been proposed to improve either utility or communication issue.

LDP Mechanisms for numerical data Similar to centralized DP, the Laplace Mechanism [19] can be adapted for local applications. Alternatively, Duchi et al. introduced a method for estimating mean of numerical values [18]. To overcome the computational and storage complexities associated with this method, an enhanced technique [17] was later developed to convert numerical inputs into binary outputs based on specific probabilities. More recently, Wang et al. [30] introduced the Piecewise Mechanism (PM) to enhance estimation accuracy, while Li et al. [26] developed the Square Wave (SW) mechanism to facilitate the estimation of numerical distributions. Subsequently, Duan et al. [15] proposed a unified framework to evaluate these mechanisms.

10 Conclusion

In this work, we study the problem of range mean estimation under Local Differential Privacy. We design a novel framework *PrivRM* for enhancing estimation accuracy by enabling mechanisms to concentrate on the target range. Two implementations are also developed, namely *PrivRM^I* and *PrivRM^{*}*, which are adaptable to all the existing LDP mechanisms for numerical value perturbation. We also provide a black-box guideline to suggest which implementation to use in different settings. Moreover, we further design a distribution-aware strategy on perturbation space adjustment, which improves estimation accuracy especially for skewed data distributions. Finally we validate the effectiveness of our methods through extensive experimental evaluations.

As for future work, we plan to extend our framework to handle more complicated range query, such as range distribution estimation. We also plan to explore efficient solutions to answer successive range queries with arbitrary query ranges.

Acknowledgments

This work was supported by the National Natural Science Foundation of China (Grant No: 62372122, 92270123), and the Research Grants Council, Hong Kong SAR, China (Grant No: 15225921, 15208923, 25207224, and C2003-23Y).

References

- [1] 2019. New york taxi trip record data. <https://www.nyc.gov/site/tlc/about/tlc-trip-record-data.page>
- [2] 2020. Frequent itemset mining dataset repository. <http://fimi.ua.ac.be/data/>
- [3] 2024. PrivRM: A Framework for Range Mean Estimation under Local Differential Privacy (Full version). <https://github.com/Yunicom-YLT/Range-mean>
- [4] 2024. USA Real Estate Dataset. <https://www.kaggle.com/datasets/ahmedshahriarsakib/usa-real-estate-dataset>
- [5] Ergute Bao, Yin Yang, Xiaokui Xiao, and Bolin Ding. 2021. CGM: an enhanced mechanism for streaming data collection with local differential privacy. *Proceedings of the VLDB Endowment* 14, 11 (2021), 2258–2270.
- [6] Raef Bassily, Kobbi Nissim, Uri Stemmer, and Abhradeep Guha Thakurta. 2017. Practical locally private heavy hitters. *Advances in Neural Information Processing Systems* 30 (2017).
- [7] Raef Bassily and Adam Smith. 2015. Local, private, efficient protocols for succinct histograms. In *Proceedings of the forty-seventh annual ACM symposium on Theory of computing*. 127–135.
- [8] Mark Bun, Jelani Nelson, and Uri Stemmer. 2019. Heavy hitters and the structure of local privacy. *ACM Transactions on Algorithms (TALG)* 15, 4 (2019), 1–40.
- [9] Graham Cormode, Somesh Jha, Tejas Kulkarni, Ninghui Li, Divesh Srivastava, and Tianhao Wang. 2018. Privacy at scale: Local differential privacy in practice. In *Proceedings of the 2018 International Conference on Management of Data*. 1655–1658.
- [10] Graham Cormode, Tejas Kulkarni, and Divesh Srivastava. 2018. Marginal release under local differential privacy. In *Proceedings of the 2018 International Conference on Management of Data*. 131–146.
- [11] Bolin Ding, Janardhan Kulkarni, and Sergey Yekhanin. 2017. Collecting telemetry data privately. *Advances in Neural Information Processing Systems* 30 (2017).
- [12] Linkang Du, Zhikun Zhang, Shaojie Bai, Changchang Liu, Shouling Ji, Peng Cheng, and Jiming Chen. 2021. AHEAD: adaptive hierarchical decomposition for range query under local differential privacy. In *Proceedings of the 2021 ACM SIGSAC Conference on Computer and Communications Security*. 1266–1288.
- [13] Rong Du, Qingqing Ye, Yue Fu, Haibo Hu, Jin Li, Chengfang Fang, and Jie Shi. 2023. Differential aggregation against general colluding attackers. In *2023 IEEE 39th International Conference on Data Engineering (ICDE)*. IEEE, 2180–2193.
- [14] Yuntao Du, Yujia Hu, Zhikun Zhang, Ziquan Fang, Lu Chen, Baihua Zheng, and Yunjun Gao. 2023. Ldptrace: Locally differentially private trajectory synthesis. *Proceedings of the VLDB Endowment* 16, 8 (2023), 1897–1909.
- [15] Jiawei Duan, Qingqing Ye, and Haibo Hu. 2022. Utility analysis and enhancement of LDP mechanisms in high-dimensional space. In *2022 IEEE 38th International Conference on Data Engineering (ICDE)*. IEEE, 407–419.
- [16] John C Duchi, Michael I Jordan, and Martin J Wainwright. 2013. Local privacy and statistical minimax rates. In *2013 IEEE 54th annual symposium on foundations of computer science*. IEEE, 429–438.
- [17] John C Duchi, Michael I Jordan, and Martin J Wainwright. 2018. Minimax optimal procedures for locally private estimation. *J. Amer. Statist. Assoc.* 113, 521 (2018), 182–201.
- [18] Duchi, John C and Jordan, Michael I and Wainwright, Martin J. 2014. Privacy aware learning. *Journal of the ACM (JACM)* 61, 6 (2014), 1–57.
- [19] Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. 2006. Calibrating noise to sensitivity in private data analysis. In *Theory of Cryptography: Third Theory of Cryptography Conference, TCC 2006, New York, NY, USA, March 4-7, 2006. Proceedings 3*. Springer, 265–284.
- [20] Úlfar Erlingsson, Vasyl Pihur, and Aleksandra Korolova. 2014. Rappor: Randomized aggregatable privacy-preserving ordinal response. In *Proceedings of the 2014 ACM SIGSAC conference on computer and communications security*. 1054–1067.
- [21] Xiaolan Gu, Ming Li, Yueqiang Cheng, Li Xiong, and Yang Cao. 2020. {PCKV}: Locally differentially private correlated {Key-Value} data collection with optimized utility. In *29th USENIX security symposium (USENIX security 20)*. 967–984.
- [22] Michael Hay, Vibhor Rastogi, Gerome Miklau, and Dan Suciu. 2010. Boosting the accuracy of differentially private histograms through consistency. *Proceedings of the VLDB Endowment* 3, 1–2 (2010), 1021–1032.
- [23] Peter Kairouz, Sewoong Oh, and Pramod Viswanath. 2014. Extremal mechanisms for local differential privacy. *Advances in neural information processing systems* 27 (2014).
- [24] Tejas Kulkarni. 2019. Answering range queries under local differential privacy. In *Proceedings of the 2019 International Conference on Management of Data*. 1832–1834.
- [25] Xiaoguang Li, Ninghui Li, Wenhai Sun, Neil Zhenqiang Gong, and Hui Li. 2023. Fine-grained poisoning attack to local differential privacy protocols for mean and variance estimation. In *32nd USENIX Security Symposium (USENIX Security 23)*. 1739–1756.
- [26] Zitao Li, Tianhao Wang, Milan Lopuhaä-Zwakenberg, Ninghui Li, and Boris Škoric. 2020. Estimating numerical distributions under local differential privacy. In *Proceedings of the 2020 ACM SIGMOD International Conference on Management of Data*. 621–635.
- [27] Alexander McFarlane Mood. 1950. Introduction to the Theory of Statistics. (1950).

- [28] Zhan Qin, Ting Yu, Yin Yang, Issa Khalil, Xiaokui Xiao, and Kui Ren. 2017. Generating synthetic decentralized social graphs with local differential privacy. In *Proceedings of the 2017 ACM SIGSAC conference on computer and communications security*. 425–438.
- [29] Abhradeep Guha Thakurta, Andrew H Vyrros, Umesh S Vaishampayan, Gaurav Kapoor, Julien Freuding, Vipul Ved Prakash, Arnaud Legendre, and Steven Duplinsky. 2017. Emoji frequency detection and deep link frequency. US Patent 9,705,908.
- [30] Ning Wang, Xiaokui Xiao, Yin Yang, Jun Zhao, Siu Cheung Hui, Hyejin Shin, Junbum Shin, and Ge Yu. 2019. Collecting and analyzing multidimensional data with local differential privacy. In *2019 IEEE 35th International Conference on Data Engineering (ICDE)*. IEEE, 638–649.
- [31] Tianhao Wang, Jeremiah Blocki, Ninghui Li, and Somesh Jha. 2017. Locally differentially private protocols for frequency estimation. In *26th USENIX Security Symposium (USENIX Security 17)*. 729–745.
- [32] Tianhao Wang, Bolin Ding, Jingren Zhou, Cheng Hong, Zhicong Huang, Ninghui Li, and Somesh Jha. 2019. Answering multi-dimensional analytical queries under local differential privacy. In *Proceedings of the 2019 International Conference on Management of Data*. 159–176.
- [33] Stanley L Warner. 1965. Randomized response: A survey technique for eliminating evasive answer bias. *J. Amer. Statist. Assoc.* 60, 309 (1965), 63–69.
- [34] Yongji Wu, Xiaoyu Cao, Jinyuan Jia, and Neil Zhenqiang Gong. 2022. Poisoning Attacks to Local Differential Privacy Protocols for {Key-Value} Data. In *31st USENIX Security Symposium (USENIX Security 22)*. 519–536.
- [35] Xiaokui Xiao, Guozhang Wang, and Johannes Gehrke. 2010. Differential privacy via wavelet transforms. *IEEE Transactions on knowledge and data engineering* 23, 8 (2010), 1200–1214.
- [36] Min Ye and Alexander Barg. 2018. Optimal schemes for discrete distribution estimation under locally differential privacy. *IEEE Transactions on Information Theory* 64, 8 (2018), 5662–5676.
- [37] Qingqing Ye, Haibo Hu, Man Ho Au, Xiaofeng Meng, and Xiaokui Xiao. 2020. LF-GDPR: A framework for estimating graph metrics with local differential privacy. *IEEE Transactions on Knowledge and Data Engineering* 34, 10 (2020), 4905–4920.
- [38] Qingqing Ye, Haibo Hu, Kai Huang, Man Ho Au, and Qiao Xue. 2023. Stateful switch: Optimized time series release with local differential privacy. In *IEEE INFOCOM 2023-IEEE Conference on Computer Communications*. IEEE, 1–10.
- [39] Qingqing Ye, Haibo Hu, Ninghui Li, Xiaofeng Meng, Huadi Zheng, and Haotian Yan. 2021. Beyond value perturbation: Local differential privacy in the temporal setting. In *IEEE INFOCOM 2021-IEEE Conference on Computer Communications*. IEEE, 1–10.
- [40] Qingqing Ye, Haibo Hu, Xiaofeng Meng, and Huadi Zheng. 2019. PrivKV: Key-value data collection with local differential privacy. In *2019 IEEE Symposium on Security and Privacy (SP)*. IEEE, 317–331.
- [41] Yuemin Zhang, Qingqing Ye, Rui Chen, Haibo Hu, and Qilong Han. 2023. Trajectory Data Collection with Local Differential Privacy. *Proceedings of the VLDB Endowment* 16, 10 (2023), 2591–2604.
- [42] Zhikun Zhang, Tianhao Wang, Ninghui Li, Shibo He, and Jiming Chen. 2018. CALM: Consistent adaptive local marginal for marginal release under local differential privacy. In *Proceedings of the 2018 ACM SIGSAC Conference on Computer and Communications Security*. 212–229.

Received October 2024; revised January 2025; accepted February 2025