

Private and Utility Enhanced Recommendations with Local Differential Privacy and Gaussian Mixture Model

Jeyamohan Neera, Xiaomin Chen, Nauman Aslam, Kezhi Wang and Zhan Shu

Abstract—Recommendation systems rely heavily on behavioural and preferential data (e.g. ratings and likes) of a user to produce accurate recommendations. However, such unethical data aggregation and analytical practices of Service Providers (SP) causes privacy concerns among users. Local differential privacy (LDP) based perturbation mechanisms address this concern by adding noise to users' data at the user-side before sending it to the SP. The SP then uses the perturbed data to perform recommendations. Although LDP protects the privacy of users from SP, it causes a substantial decline in recommendation accuracy. We propose an LDP-based Matrix Factorization (MF) with a Gaussian Mixture Model (MoG) to address this problem. The LDP perturbation mechanism, i.e., Bounded Laplace (BLP), regulates the effect of noise by confining the perturbed ratings to a predetermined domain. We derive a sufficient condition of the scale parameter for BLP to satisfy ϵ -LDP. We use the MoG model at the SP to estimate the noise added locally to the ratings and the MF algorithm to predict missing ratings. Our LDP based recommendation system improves the predictive accuracy without violating LDP principles. We demonstrate that our method offers a substantial increase in recommendation accuracy under a strong privacy guarantee through empirical evaluations on three real-world datasets, i.e., Movielens, Libimseti and Jester.

Index Terms—Data Privacy, Gaussian Mixture Model, Local Differential Privacy, Recommendation Systems

1 INTRODUCTION

THE proliferation of smartphones is boosting the usage of online shopping platforms. With more and more retailers moving online, users feel overwhelmed with too many options and have trouble finding a product or service to fulfil their expectations. Most online shopping platforms use recommendation systems so that users can find items that could interest them.

Collaborative Filtering (CF) is a recommendation model widely used to predict users' preference for unpurchased items. Although CF offers higher recommendation accuracy, it causes privacy risks as service providers (SPs) gather a large amount of user data to predict purchasing behaviour patterns of users. Narayanan *et al.* [1] demonstrated how analyzing users' historical ratings can disclose sensitive information such as users' political preference, health-related information and sometimes even their sexual orientation. Hence, it is pivotal for SPs to protect the privacy of users while providing suitable personalized recommendations.

Differential Privacy (DP) is a popular tool that can guarantee strong privacy protection even when the adversary owns a considerable amount of auxiliary information about the user [30]. Most of the existing works on DP based privacy protection methods focused on protecting the privacy of users against a third-party adversary and assume that the risk of the SP causing privacy violation is minimal. Unfortunately, many SPs are apt to gather more data from

users than they need and continue to procure sensitive information about users' behaviour for their own added benefits. Cambridge Analytica investigations [27] reveal the dangerous consequences of such harmful and unethical data aggregation practices. Not only could untrustworthy SPs violate users' privacy, but even trustworthy SPs might encounter an accidental privacy leakage as they own an enormous amount of sensitive user information. Narayanan *et al.* [1] deanonymize the Netflix rating data to show how trusted SPs could cause accidental data leakage.

Local Differential Privacy (LDP) [32] has attracted much attention as it can provide a strong privacy guarantee in a setting where SPs are untrustworthy. Many researchers [4], [5], [6] have adopted LDP to protect the privacy of users in recommendation systems. Each user adds noise to their data locally in LDP-based privacy protection models and forwards the perturbed data to the SP. As the original data never leaves the user device, users are guaranteed plausible deniability. Adopting LDP in recommendation systems cause low data utility for SPs. The recommendation accuracy is comparatively higher in DP-based recommendation systems as they perturb a query output, whereas LDP-based recommendation systems add noise to the individual data point. Therefore, it is crucial to design LDP based recommendation system that provides strong privacy protections to users and simultaneously offers higher data utility to the SP.

Motivated by this, we propose an LDP-based recommendation system that perturbs original ratings of a user within a predefined domain using the Bounded Laplace (BLP) mechanism. We then use a Mixture of Gaussian (MoG) model to estimate the aggregate noise at the SP to enhance the data utility. The main contributions of this work are

- Jeyamohan Neera, Xiaomin Chen, Nauman Aslam and Kezhi Wang are with Northumbria University, UK. E-mail: jeyamohan.neera@northumbria.ac.uk, xiaomin.chen@northumbria.ac.uk, nauman.aslam@northumbria.ac.uk, kezhi.wang@northumbria.ac.uk
- Zhan Shu is with University of Alberta, Canada. E-mail: zshu1@ualberta.ca

Manuscript received ; revised

listed below:

- We introduce BLP as the input rating perturbation mechanism to increase the recommendation accuracy of the LDP-based recommendation systems. To the best of our knowledge, we are the first to introduce BLP as the input data perturbation mechanism and to provide a sufficient condition for BLP to satisfy ϵ -local differential privacy in recommendation systems. We empirically evaluate the role the BLP mechanism plays in enhancing predictive accuracy in Section 5.3.2.
- The probability density function of Bounded Laplace noise is conditional on the input rating matrix, unlike the Laplace mechanism, where there are no bounding constraints to restrict the noise samples. Hence we perform a theoretical analysis to identify and yield a noise distribution for the Bounded Laplace mechanism. We derive a closed-form probability density function for noise drawn from BLP for a given dataset.
- We introduce a noise estimation component at SP to further increase the predictive accuracy of the recommendation system. Perturbation of each user's rating leads to higher predictive error, which increases linearly with the number of users and items. We adopt Matrix Factorization (MF) with a Mixture of Gaussian (MoG) to estimate the aggregated noise at SP and at the same time to predict missing ratings. This novel approach tackles data utility issues found in LDP based recommendation systems. We empirically evaluate the effect of MF with MoG in terms of achieving higher recommendation accuracy in Section 5.3.3. We also show that the proposed LDP based recommendation model outperforms the existing LDP based recommendation models such as [4] and [6] in Section 5.3.4.
- Our approach causes much lower communication cost compared to existing LDP based recommendation systems e.g. [4]. Users only need to transmit each perturbed rating once to the SP in our proposed method. On the contrary, in other systems such as [4], the information exchange between a user and the SP continues for several iterations until the solution converges.

Our method protects users' privacy from untrustworthy service providers. However, there are some cautions and limitations that we need to indicate. Firstly, we assume that each user sends a single rating to the service provider at any given time, and this rating is independent of other users' ratings. We presume all the ratings are sensitive and essential in building a behavioural profile for a user. Hence we perturb all the ratings of the users that cause heavy utility loss. We provide a solution to address this issue - combining a noise estimation model (MoG) with the recommendation algorithm to sanitise the obfuscated data at the service provider. Secondly, we only mask the rating scores of users to items, but not the set of items a user has rated. The exposure of user-item association can harm user privacy to some extent. However, hiding the relationship between users and items will further reduce data utility. The recently

proposed methods such as federated learning-based models [39] [40] require a significant change to the system architecture. On the contrary, we can apply our proposed privacy-enhanced recommendation system to an existing centralised recommendation model with satisfactory recommendation accuracy and low communication and computation costs. We will seek solutions to strengthen privacy protections further against the exposure of user-item linkage in future work.

In Table 1, we list notations frequently used in this paper.

TABLE 1
Notations

Notation	Meaning
Pr	Probability
R	Original rating matrix
R^*	Perturbed rating matrix
N	BLP noise matrix
r_{ij} or r	Original rating
r_{ij}^* or r^*	Perturbed rating
n_{ij} or n	BLP noise
l	Minimum value in rating scale
u	Maximum value in rating scale
U	User Latent Factor Matrix
V	Item Latent Factor Matrix
u_i	Latent factors of user i in latent matrix U
v_j	Latent factors of item j in latent matrix V

2 RELATED WORK

Generally, recommendation systems use privacy protection models based on techniques such as obfuscation [28] and perturbation [29]. These approaches introduce random noise to data. Yet, the magnitude of noise added using these methods cannot be calibrated easily. DP is another popular perturbation approach used in privacy protection models. DP based methods are proven to be a stronger solution for privacy protection in various applications compared to other perturbation methods. DP provides an information-theoretic guarantee of strong privacy protection regardless of how much knowledge the adversary possesses. Unlike other perturbation approaches, in DP, the calibration of noise depends on the sensitivity of the query and the level of privacy offered to the user.

Many DP based recommendation systems assume a trusted SP who collects users' ratings and releases information related to users' preferences under a differential privacy guarantee. McSherry and Mironov [2] are the first ones to integrate the DP based privacy protection model with collaborative filtering-based recommendation systems. In their method, SP is considered trustworthy. They build a covariance matrix using the user's original ratings that resemble the similarity between users. They then apply the Laplace mechanism to perturb the covariance matrix before predicting missing ratings. In their method, the trusted SP still has access to sensitive unperturbed data of a user. Yakut and Polat [34] also introduce a DP based recommendation system where user's original ratings are stored at SP using a perturbation approach which provides uncertainty over user's actual ratings. This method also ensures that some user profiles contain fake ratings depending on the privacy budget set by the SP. Even though DP based recommendation models offer privacy protection to users from third

party adversaries, they enable SP to collect the original ratings from users, which in return causes privacy concerns.

Hence, the attention of researchers is gradually shifting from DP to LDP. Many applications have adopted LDP to deal with untrustworthy SPs. Erlingsson *et al.* [7] propose the first local differentially private perturbation mechanism, RAPPOR, for crowd-sourcing data aggregation. Google uses this mechanism [7] to collect users' Chrome usage statistics privately. RAPPOR uses a randomized response (RR) mechanism to perturb each bit of the data independently before sending it to Google for further analysis. However, high communication overhead during the data collection phase is a drawback found in this method. Bassily *et al.* [35] propose another perturbation mechanism to address the issue found in RAPPOR. Each user report one randomly chosen bit rather than reporting d number of bits back to the SP using the succinct histogram (SH) mechanism. However, the SH perturbation mechanism is more suitable to the simple numeric or categorical attributes, and it is not appropriate for more complex data mining tasks.

Qin *et al.* [36] also propose an algorithm for performing heavy hitters estimation under the guarantee of LDP. In another work, Wang *et al.* [37] propose another LDP algorithm for data aggregation and decoding. These two methods can generate perturbed data and reconstruct the statistical characteristics. However, they cannot recreate the cross-correlation relationship between data. To address this problem, Zhang *et al.* [38] propose an LDP algorithm that can reconstruct cross-correlation relationships among high-dimensional data. In their method, marginal tables are generated and then perturbed. They send only the noisy marginal tables to the SP. However, this method is not suitable for collaborative filtering based recommendation systems as these systems require users to send their historical data.

Several works have investigated using LDP in CF-based recommendation systems. In their work, Liu *et al.* [31] propose a privacy-preserving recommendation system that uses a randomized perturbation mechanism. First, noise is added to users' ratings locally on the user-side through a randomized perturbation method. Additionally, they also add noise to the correlation computation method at SP. Even though this method can guarantee more privacy protection, it incurs more predictive accuracy loss. Meng *et al.* [9] address this problem by concentrating on ratings that are considered sensitive. They divide a user's historical ratings into sensitive and non-sensitive ratings. They use a large magnitude of noise to perturb sensitive ratings. Hence, sensitive ratings receive better privacy protection, while the recommendation system can achieve improved predictive accuracy. However, the distinction between sensitive and non-sensitive data can vary from user to user and cannot be generalized to all users.

Shen and Jin [33] propose an instance-based relaxed admissible mechanism to perturb users' private data. They aim to hide users' preference towards an item from an untrusted data aggregator. However, this method can still reveal users' preferences towards an item category. Hua *et al.* [5] propose another LDP based recommendation model where the SP uses LDP based MF method to compute item profile latent factors. Subsequently, SP sends these item profiles latent

factors to the users for computation of user latent factors. Each user then sends the updated item latent factors back to the SP. This method requires the users to remain online during the whole MF process. Their proposed model adds additional communication and processing cost on the user side.

Shin *et al.* [4] also propose an LDP-based recommendation model where they use a randomized response mechanism to perturb data on the user side. In their method, instead of sending the item latent factors back to the SP, each user sends back the perturbed gradient of their user latent factor. This method incurs additional processing and communication overhead to the user-side as same as [5]. In another work, Berlioz *et al.* [6] have investigated the effect of rating perturbation in different stages of the recommendation process. They evaluate the role input and output perturbation mechanisms play on predictive accuracy.

To summarize, existing LDP based recommendation systems suffer from low predictive accuracy and communication/computational overhead. Current works do not perform any analysis on the perturbed data to eliminate the noise caused by randomization. Additionally, some of the proposed methods focus on frequency estimation, such as heavy hitter identification. Hence these methods are not suitable for collaborative filtering as they focus only on a specific candidate set and cannot identify the correlation between users/items. Therefore, we propose an LDP based recommendation system, which perturbs users' ratings on the user-side and estimates the noises added on the SP side to ensure high predictive accuracy.

3 PRELIMINARIES

3.1 Differential Privacy

DP-based privacy protection models are relevant in settings where the SP is trusted and aggregates users' original data. Assume two adjacent data sets D and D' where D' differs from D by one record.

Definition 1. A randomized mechanism M satisfies ϵ -differential privacy if for any adjacent datasets D and D' , and any subset S of all possible outputs, we have the following inequality:

$$Pr[M(D) \in S] \leq e^\epsilon \times Pr[M(D') \in S],$$

where ϵ is the privacy budget.

DP limits an adversary from inferring whether an input data set D or D' produced the given output S and the privacy budget ϵ controls the privacy loss. The smaller the value of the privacy budget ϵ , the lower the confidence the adversary has in distinguishing whether dataset D or D' produced the output. Hence, DP provides a higher degree of privacy protection for lower values of privacy budget ϵ .

Definition 2. Given a query $f : D \rightarrow \mathbb{R}$, the sensitivity of f , Δf , can be defined as:

$$\Delta f = \max_{D, D'} \|f(D) - f(D')\|.$$

The sensitivity of a function indicates how much noise is required to perturb a query result. It parametrizes the maximum difference a single record can make on the output of a query.

3.1.1 Post Processing

Dwork *et al.* [30] proved that a differentially private result is immune to any privacy attacks even after post-processing. This property is key to our work, as SP can carry out any computation on aggregates of differentially private output without violating the principles of DP. The following definition clearly defines this property.

Definition 3. Let M be a randomized algorithm that is ϵ -differentially private. Let g be an arbitrary function. Then, $g \circ M$ is also ϵ -differentially private."

3.2 Local Differential Privacy

Definition 4. A randomized mechanism M satisfies ϵ -LDP if for all possible pairs of user input x, x' and any subset y of all possible outcomes, we have the following inequality:

$$Pr[M(x) \in y] \leq e^\epsilon \times Pr[M(x') \in y].$$

In the LDP setting, the data of each user is perturbed locally before being sent to the SP. So the SP aggregates the perturbed data instead of the original data. Therefore, SP cannot infer any information about the actual data by observing the perturbed output even though they possess substantial background knowledge about the user. In this regard, LDP offers plausible deniability to users. Intuitively, LDP ensures that the SP cannot infer whether a user's input x or x' produce the output y with confidence.

3.3 Laplace mechanism

The Laplace mechanism adds random noise drawn from Laplace distribution to ensure ϵ -differential privacy. We use the notation $Lap(0, b_{lap})$ to indicate that the Laplace mechanism uses a Laplace distribution with mean 0 and scale parameter (i.e. variance) b_{lap} to sample noise.

Definition 5. Given a query $f : D \rightarrow \mathbb{R}$, the randomized mechanism M satisfies ϵ -differential privacy if:

$$M(D) = f(D) + Lap(0, b_{lap}).$$

The scale parameter b_{lap} controls the width of Laplace distribution. If Δf is the sensitivity of the query f and ϵ is the privacy budget, then the scale parameter b_{lap} of Laplace distribution can be determined as:

$$b_{lap} = \frac{\Delta f}{\epsilon}.$$

Hence, the width of the Laplace distribution is dependent on sensitivity Δf and privacy budget ϵ .

3.4 Matrix Factorization

Matrix Factorization algorithm is the state-of-the-art technology used in CF-based recommendation systems. Many E-commerce platforms prefer MF over other methods due to its higher predictive accuracy and computational scalability. The rating matrix R that contains ratings of m users over n items acts as input for the MF algorithm. Each element r_{ij} in the rating matrix indicates the rating of user $i \in \{1, 2, \dots, m\}$ on item $j \in \{1, 2, \dots, n\}$. MF algorithm predicts the missing rating by modelling the interactions between users and items as the inner product of latent factor spaces. MF algorithm factorizes the given rating matrix

R into two latent matrices: U (user latent factor matrix) and V (item latent factor matrix). MF obtains the user and item latent matrices by minimizing the squared error for all known ratings in the rating matrix.

$$\min_{U, V} \sum_{r_{ij} \in R} [r_{ij} - u_i^T v_j]^2. \quad (1)$$

In Eq. (1), u_i represents the relationship between user i and the latent factors in the user latent matrix U . Similarly, v_j represents the relationship between item j and the latent factors in the item latent matrix V . The non-convex optimization problem given by Eq. (1) is solved using either stochastic gradient descent (SGD) or alternating least squares (ALS) method. After obtaining the latent matrices, MF predicts the missing rating of a user on an item using the dot product of the corresponding user and item latent column vectors:

$$\hat{r}_{ij} = u_i^T v_j.$$

4 LOCAL DIFFERENTIAL PRIVACY RECOMMENDATION WITH BLP AND MoG

Our proposed recommendation model is applicable in a setting where the users are cautious about sharing sensitive information with an untrustworthy SP. Fig. 1 illustrates the proposed recommendation system. An LDP mechanism, BLP, perturbs the actual ratings of a user before sending to the SP. Hence, the SP can only aggregate perturbed ratings from the users. At the SP, MF with MoG model estimates the noise added to the ratings and perform missing rating prediction. The post-processing property of LDP implies that further processing a perturbed output of a ϵ -differentially private mechanisms does not cause any adverse effects on privacy protection [30]. Since LDP mechanisms are immune to post-processing, estimating noise at the SP-side does not cause any additional privacy risk to users. We will describe each component of the system in detail in this section.

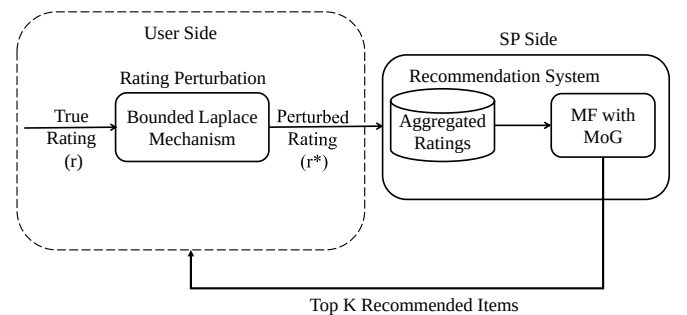


Fig. 1. LDP based MF recommendation with MoG

4.1 LDP Rating Perturbation

4.1.1 Bounded Laplace Mechanism

As discussed in section 3, the Laplace mechanism achieves ϵ -differential privacy by sampling random noise within the range of $-\infty$ to ∞ . The perturbed output thus falls within the domain of $-\infty$ to ∞ . For example, the Laplace

mechanism might produce a negative result as a perturbed output while perturbing a user rating. Although this negative output holds no physical meaning in terms of the rating scale, it is still a valid output of the Laplace mechanism. Such inconsistent perturbed ratings have an immense effect on the predictive accuracy of MF based recommendation systems.

We use BLP as an input perturbation mechanism to increase the predictive accuracy of LDP based recommendation systems. The BLP mechanism ignores off-limit values and samples noise for a given input rating continuously until a perturbed rating falls within the predefined output domain. Given an input rating r , the BLP mechanism continuously samples noise from a Laplace distribution until the perturbed rating r^* falls within the predefined output domain, i.e. $l \leq r^* \leq u$, where l is the minimum and, u is the maximum value of the given rating scale.

Bounded Laplace mechanism can be defined using the probability density function (pdf) as below [26]:

Definition 6. Given a domain interval of (l, u) , input $r \in [l, u]$ and the scale parameter $b > 0$, the Bounded Laplace mechanism W , is given by the conditional probability density function :

$$f_W(r^*) = \begin{cases} \frac{1}{C(r)} \frac{1}{2b} e^{-\frac{|r^*-r|}{b}}, & \text{if } r^* \in [l, u], \\ 0, & \text{if } r^* \notin [l, u], \end{cases}$$

where $C(r) = \int_l^u \frac{1}{2b} e^{-\frac{|r^*-r|}{b}} dr^*$ is a normalization factor dependent on input r .

Lemma 1. The normalization factor $C(r)$ is given by:

$$C(r) = 1 - \frac{1}{2} \left(e^{-\frac{r-l}{b}} + e^{-\frac{u-r}{b}} \right).$$

Proof.

$$\begin{aligned} C(r) &= \int_l^u \frac{1}{2b} e^{-\frac{|r^*-r|}{b}} dr^* \\ &= \int_l^r \frac{1}{2b} e^{\frac{r^*-r}{b}} dr^* + \int_r^u \frac{1}{2b} e^{-\frac{r^*-r}{b}} dr^* \\ &= \frac{b}{2b} \left[e^{\frac{r^*-r}{b}} \right]_l^r + \frac{b}{2b} \left[-e^{-\frac{r^*-r}{b}} \right]_r^u \\ &= 1 - \frac{1}{2} \left(e^{-\frac{r-l}{b}} + e^{-\frac{u-r}{b}} \right). \end{aligned}$$

Assume that r and r' are a pair of possible inputs to a randomized mechanism and $r' = r + z$. We define $F(r, z)$ as:

$$F(r, z) = \frac{C(r+z)}{C(r)} e^{\frac{|r'-r|}{b}}.$$

Lemma 2. Let $0 \leq z \leq \Delta f$, then,

$$\max_{\substack{r, r' \in [l, u] \\ 0 \leq z \leq \Delta f}} F(r, z) = \frac{C(l + \Delta f)}{C(l)} e^{\frac{\Delta f}{b}}.$$

Proof. The full proof is given in Appendix A. $\square$

We define ΔC for later use:

$$\Delta C = \frac{C(l + \Delta f)}{C(l)}.$$

Theorem 1. When scale parameter $b \geq \frac{\Delta f}{\varepsilon - \log \Delta C}$, it is sufficient to show that the Bounded Laplace mechanism W satisfies ε -local differential privacy

Proof. Assume that r and r' are a pair of possible inputs to a Bounded Laplace mechanism and $r' = r + z$. Let $0 \leq z \leq \Delta f$. r^* represents a perturbed output produced by the BLP mechanism. Given the domain of the perturbed output is $[l, u]$, we can note that,

$$Pr(W(r) \in [l, u]) = \frac{1}{C(r)} Pr(M(r) \in [l, u]),$$

where M represents the Laplace mechanism.

We aim to find a condition under which W satisfies ε -local differential privacy. Based on the LDP definition, we can note that,

$$\begin{aligned} Pr(W(r) \in [l, u]) &\leq e^\varepsilon Pr(W(r') \in [l, u]), \\ \frac{1}{C(r)} Pr(M(r) \in [l, u]) &\leq e^\varepsilon \frac{1}{C(r')} Pr(M(r') \in [l, u]). \end{aligned}$$

Given that $Pr(M(r) \in [l, u]) = \int_l^u \frac{1}{2b} e^{-\frac{|r^*-r|}{b}} dr^*$, we have,

$$\frac{1}{C(r)} \int_l^u \frac{e^{-\frac{|r^*-r|}{b}}}{2b} dr^* \leq e^\varepsilon \frac{1}{C(r')} \int_l^u \frac{e^{-\frac{|r^*-r'|}{b}}}{2b} dr^*. \quad (2)$$

A lower bound for $e^\varepsilon \frac{1}{C(r')} \int_l^u \frac{e^{-\frac{|r^*-r'|}{b}}}{2b} dr^*$ can be obtained using the triangle inequality, i.e.

$$\begin{aligned} |r^* - r'| &\leq |r^* - r| + |r' - r|, \\ e^\varepsilon \frac{1}{C(r')} \int_l^u \frac{e^{-\frac{|r^*-r'|}{b}}}{2b} dr^* &\geq e^\varepsilon \frac{1}{C(r')} \int_l^u \frac{e^{-\frac{|r^*-r|+|r'-r|}{b}}}{2b} dr^*, \\ e^\varepsilon \frac{1}{C(r')} \int_l^u \frac{e^{-\frac{|r^*-r'|}{b}}}{2b} dr^* &\geq e^{\varepsilon - \frac{|r'-r|}{b}} \frac{1}{C(r')} \int_l^u \frac{e^{-\frac{|r^*-r|}{b}}}{2b} dr^*. \end{aligned}$$

To ensure Eq. (2) hold, it is sufficient to show that:

$$\frac{1}{C(r)} \int_l^u \frac{1}{2b} e^{-\frac{|r^*-r|}{b}} dr^* \leq e^{\varepsilon - \frac{|r'-r|}{b}} \frac{1}{C(r')} \int_l^u \frac{1}{2b} e^{-\frac{|r^*-r|}{b}} dr^*. \quad (3)$$

The inequality given by Eq. (3) can be further reduced as,

$$\frac{C(r)}{C(r')} e^{\varepsilon - \frac{|r'-r|}{b}} \geq 1.$$

From Lemma 2 we can note that,

$$\frac{C(r')}{C(r)} e^{\frac{|r'-r|}{b}} \leq \Delta C e^{\frac{\Delta f}{b}},$$

Equivalently,

$$\frac{C(r)}{C(r')} e^{\varepsilon - \frac{|r'-r|}{b}} \geq \frac{1}{\Delta C} e^{\varepsilon - \frac{\Delta f}{b}}.$$

We find a lower bound for $\frac{C(r)}{C(r')} e^{\varepsilon - \frac{|r'-r|}{b}}$ and proceed to find a condition for Eq. (3) to hold.

To make Eq. (3) hold, it is sufficient to show that,

$$1 \leq e^{\varepsilon - \frac{\Delta f}{b}} \frac{1}{\Delta C}.$$

or equivalently,

$$b \geq \frac{\Delta f}{\varepsilon - \log(\Delta C)}.$$

Theorem 1 provides the scale parameter for BLP to satisfy ϵ -local differential privacy. It also demonstrates that BLP cannot satisfy ϵ -local differential privacy when inheriting the scale parameter from the Laplace mechanism. In our recommendation system, we use BLP as an input rating perturbation mechanism. Input perturbation mechanism calibrates the magnitude of noise added to original ratings according to the sensitivity given by $\Delta f = u - l$.

We define ΔC as:

$$\begin{aligned}\Delta C &= \frac{C(l + \Delta f)}{C(l)} \\ &= \frac{1 - \frac{1}{2}(e^{-\frac{\Delta f}{b}} + e^{-\frac{u - \Delta f - l}{b}})}{1 - \frac{1}{2}(1 + e^{-\frac{u - l}{b}})}.\end{aligned}$$

When $\Delta f = u - l$,

$$\Delta C = \frac{1 - \frac{1}{2}(1 + e^{-\frac{(u-l)}{b}})}{1 - \frac{1}{2}(1 + e^{-\frac{(u-l)}{b}})} = 1.$$

Thus $\log \Delta C = 0$.

Therefore we can conclude that a sufficient condition needed for BLP mechanism to satisfy ϵ -local differential privacy in our recommendation system can be given by:

$$b \geq \frac{u - l}{\epsilon}.$$

Algorithm 1 details the stages involved in generating a perturbed rating using BLP mechanism.

Algorithm 1 BLP Mechanism for Noise Sampling

- 1: **Input to the Mechanism: Original Rating (r)**
 - 2: **Output of the Mechanism: Perturbed Rating (r^*)**
 - 3: A noise value is generated from the Laplace distribution with mean 0 and variance of b - $Lap(0, b)$
 - 4: Add noise to original rating to obtain perturbed rating:
 $r^* = r + Lap(0, b)$
 - 5: **If** ($r^* \in (l, u)$):
 - 6: Perturbed rating is set to r^*
 - 7: **else**
 - 8: repeat Step 3 until ($r^* \in (l, u)$)
 - 9: **Return** Perturbed rating to SP
-

4.1.2 BLP Noise Distribution

The noise distribution of BLP mechanism can be theoretically derived for any given dataset of true input ratings. Consider a discrete rating system containing h evenly distributed discrete ranks with step size c . The rank set is denoted by $\mathcal{Q} = \{Q_1, \dots, Q_{h-1}, Q_h\}$, and $|Q_{i+1} - Q_i| = c, 1 \leq i \leq h-1$. Let r be a true rating, its corresponding perturbed rating is $r^* = r + n$, where n is the random noise drawn by BLP mechanism. Since r^* can only take values in the set $\mathcal{Q}$, i.e. $Q_1 \leq r^* \leq Q_h$, we have the noise range for input rating r as $Q_1 - r \leq n \leq Q_h - r$. Define the probability θ_r as:

$$\theta_r = Pr(Q_1 - r \leq n < Q_h - r),$$

θ_r represents the probability that the noise variable n falls into the interval $(Q_1 - r, Q_h - r)$ given an input rating r .

- The input rating r takes values in a finite set $\mathcal{Q}$. θ_r can thus be expanded as

$$\theta_r = \sum_{Q_i \in \mathcal{Q}} Pr(r = Q_i) Pr(Q_1 - Q_i \leq n < Q_h - Q_i | r = Q_i)$$

where $Pr(r = Q_i)$ is the probability that the input rating equals to Q_i , and $Pr(Q_1 - Q_i \leq n < Q_h - Q_i | r = Q_i)$ is the conditional probability that the BLP noise lies within the interval $(Q_1 - Q_i \leq n < Q_h - Q_i)$ under the condition that the input rating is Q_i . Note that not all the input ratings in $\mathcal{Q}$ leads to the noise n falling within this particular range. When the perturbed rating $r^* \notin \mathcal{Q}$, the conditional probability $Pr(Q_1 - Q_i \leq n < Q_h - Q_i | r = Q_i)$ yields 0.

The noise added by the BLP mechanism over all possible input ratings in $\mathcal{Q}$ is a random variable ranging within $(Q_1 - Q_h \leq n < Q_h - Q_1)$. We will then divide the range into equal intervals. The length of each interval is the rank step size of c . The probability that the noise variable lies within each interval is given as:

$$\begin{aligned}Pr(Q_1 - Q_{h-t} \leq n < Q_1 - Q_{h-t-1}) &= \sum_{i=h-t}^h Pr(r = Q_i) \cdot \\ Pr(Q_1 - Q_{h-t} \leq n < Q_1 - Q_{h-t-1} | r = Q_i) \\ \forall t &= 0, \dots, h-2.\end{aligned}\tag{4}$$

and

$$\begin{aligned}Pr(Q_h - Q_{t+1} \leq n < Q_h - Q_t) &= \sum_{i=1}^t Pr(r = Q_i) \cdot \\ Pr(Q_h - Q_{t+1} \leq n < Q_h - Q_t | r = Q_i) \\ \forall t &= 2, \dots, h-1.\end{aligned}\tag{5}$$

The conditional probability is given by

$$\begin{aligned}Pr(n \in [Q_1 - Q_{h-t}, Q_1 - Q_{h-t-1}] | r = Q_i) \\ = \int_{Q_1 - Q_{h-t}}^{Q_1 - Q_{h-t-1}} \frac{1}{C_r} \frac{1}{2b} e^{-\frac{|r^* - x|}{b}} dr^* \\ = \frac{1}{C_r} \frac{1}{2b} (e^{Q_1 - Q_{h-t-1}} - e^{Q_1 - Q_{h-t}}) \\ \text{for } t = 0, \dots, h-2.\end{aligned}$$

and

$$\begin{aligned}Pr(n \in [Q_h - Q_{t+1}, Q_h - Q_t] | r = Q_i) \\ = \int_{Q_h - Q_{t+1}}^{Q_h - Q_t} \frac{1}{C_r} \frac{1}{2b} e^{-\frac{|r^* - x|}{b}} dr^* \\ = \frac{1}{C_r} \frac{1}{2b} (e^{Q_h - Q_t} - e^{Q_h - Q_{t+1}}) \\ \text{for } t = 2, \dots, h-1.\end{aligned}$$

$$\text{where } C_r = 1 - \frac{1}{2} \left(e^{-\frac{Q_i - Q_1}{b}} + e^{-\frac{Q_h - Q_i}{b}} \right).$$

4.2 Noise Estimation with Mixture of Gaussians

The MoG model is widely used to approximate probability distributions with no closed-form expression. In image processing this model is being used for the purpose of image segmentation [16], image compression [14] and background subtraction [15]. We propose an MoG with MF recommendation model to estimate the noise added to the true ratings and predict missing ratings. Since a multivariate Gaussian distribution can model the uncertainty of a noise data point, MoG is a good solution for noise estimation.

Since we add BLP noise to each true rating in the true rating matrix R , the perturbed rating matrix R^* can thus be given by:

$$R^* = R + N.$$

We aim to find a mixture of K Gaussian components which best represents the noise distribution. We assume that each noise data point n_{ij} in N is drawn from a Gaussian distribution $\mathcal{N}(n_{ij} | 0, \sigma_k^2)$ where σ_k is the standard deviation of the k -th Gaussian component. The mixture of K Gaussian components representing the noise data point n_{ij} can thus be given by:

$$p(n_{ij} | \Pi, \Sigma) \sim \sum_{k=1}^K \pi_k \mathcal{N}(n_{ij} | 0, \sigma_k^2),$$

in which π_k ($\sum_{k=1}^K \pi_k = 1$) is the mixture proportion representing the probability that n_{ij} is drawn from the k -th mixture component. $\Pi = (\pi_1, \pi_2, \dots, \pi_K)$ and $\Sigma = (\sigma_1, \sigma_2, \dots, \sigma_K)$. As given in preliminaries, each known rating in original rating matrix R can be approximated using MF as:

$$r_{ij} = (u_i^T)v_j.$$

Hence each rating r_{ij}^* in the perturbed rating matrix can be given by:

$$r_{ij}^* = r_{ij} + n_{ij} = (u_i^T)v_j + n_{ij}.$$

Subsequently, the probability distribution of perturbed rating r_{ij}^* can then be given by:

$$p(r_{ij}^* | u_i, v_j, \Pi, \Sigma) = \sum_{k=1}^K \pi_k \mathcal{N}(r_{ij}^* | (u_i^T)v_j, \sigma_k^2).$$

The likelihood of R^* can thus be given by:

$$p(R^* | V, U, \Sigma, \Pi) = \prod_{i,j \in \Omega} \sum_{k=1}^K \pi_k \mathcal{N}(r_{ij}^* | (u_i^T)v_j, \sigma_k^2),$$

where Ω represents the set of non-missing data points in perturbed rating matrix R^* . Given the likelihood, next, we derive the maximum likelihood estimates of the model parameters V, U, Σ and Π for the perturbed rating matrix R^* , i.e.:

$$\begin{aligned} \max_{V, U, \Sigma, \Pi} \mathcal{L}(R^* | V, U, \Sigma, \Pi) \\ = \sum_{i,j \in \Omega} \log \sum_{k=1}^K \left(\pi_k \mathcal{N}(r_{ij}^* | (u_i^T)v_j, \sigma_k^2) \right). \end{aligned} \quad (6)$$

The log-likelihood can be simplified as [17]:

$$\max_{U, V, \Pi, \Sigma} \sum_{i,j \in \Omega} \sum_{k=1}^K \gamma_{ijk} \left(\log \pi_k - \log \sqrt{2\pi} \sigma_k - \frac{(r_{ij}^* - (u_i^T)v_j)^2}{2\sigma_k^2} \right). \quad (7)$$

4.3 Expectation Maximization for MoG

We use Expectation-Maximization (EM) method [17] to evaluate and compute model parameters V, U, Σ and Π to maximize the likelihood function given by Eq. (6). The EM is an iterative algorithm that can be summarized as follow:

- Initialize the model parameters
- Evaluate the initial value of log-likelihood
- Expectation (E-Step) : Evaluate the posterior responsibilities using the current model parameters
- Maximization (M-Step) : Re-estimate the model parameters using the current posterior responsibilities

EM algorithm updates the parameters and alternates E-step and M-step until convergence. The standard EM algorithm estimates the mean of each cluster at every iteration. In our system, the clusters share the same parameters U and V .

At first, we randomly initialize the model parameters V, U, Σ and Π to estimate posterior responsibilities of K Gaussian components. In E-step, we estimate the posterior responsibility for each noise point n_{ij} using the current model parameters V, U, Σ and Π as:

$$\gamma_{ijk} = \frac{\pi_k \mathcal{N}(r_{ij}^* | (u_i^T)v_j, \sigma_k^2)}{\sum_{k=1}^K \pi_k \mathcal{N}(r_{ij}^* | (u_i^T)v_j, \sigma_k^2)}. \quad (8)$$

The posterior responsibility reflects the probability that k -th Gaussian component produces the noise point n_{ij} . Then in M-step, we re-estimate each model parameter V, U, Σ and Π based on the posterior responsibilities γ_{ijk} from E-step. We first update Π and Σ :

$$S_k^{(x+1)} = \sum_{i,j \in \Omega} \gamma_{ijk}^{(x)},$$

$$\pi_k^{(x+1)} = \frac{S_k^{(x+1)}}{S},$$

$$\sigma_k^2 = \frac{1}{S_k^{(x+1)}} \sum_{i,j \in \Omega} \gamma_{ijk}^{(x)} (r_{ij}^* - (u_i^T)v_j)^2, \quad (9)$$

where S is the total number of non-missing data points, $S_k^{(x+1)}$ is the sum of γ_{ijk} for k -th Gaussian component and x is the total number of iterations EM algorithm runs until convergence. Then we update the model parameters U and V . We can rewrite the portion in Eq. (7) which is related to U and V as:

$$\begin{aligned} \max_{V, U} \sum_{i,j \in \Omega} \sum_{k=1}^K \gamma_{ijk} \left(- \frac{(r_{ij}^* - (u_i^T)v_j)^2}{2\sigma_k^2} \right) \\ = - \sum_{i,j \in \Omega} \left(\sum_{k=1}^K \frac{\gamma_{ijk}}{2\sigma_k^2} \right) (r_{ij}^* - u_i^T v_j)^2 \\ = - \sum_{i,j \in \Omega} w_{ij} (r_{ij}^* - u_i^T v_j)^2, \end{aligned} \quad (10)$$

where w_{ij} represents the weight for each true rating r_{ij} , given by:

$$w_{ij} = \begin{cases} \sqrt{\sum_{k=1}^K \frac{\gamma_{ijk}}{2\sigma_k^2}}, & \text{if } i, j \in \Omega \\ 0, & \text{if } i, j \notin \Omega. \end{cases}$$

Eq. (10) is equivalent to a weighted low-rank MF problem as given below:

$$\min_{U,V} W \odot (X - UV^T)^2.$$

The weighted low-rank MF problems can be solved using methods such as Weighted Low-Rank Approximation [20], Damped Newton [19] and Weighted PCA [18]. We use Weighted PCA in this work to re-estimate model parameters U and V . The EM algorithm stops alternating between E-step and M-step when two consecutive user latent factor matrices U cause a change smaller than the given threshold value or the number of iterations reaches the pre-defined threshold. Algorithm 2 details the process of how MoG with the MF model estimates noise and predict missing ratings.

Algorithm 2 Noise Estimation and Rating Prediction Model

- 1: **Input:** Perturbed Ratings (R^*)
- 2: **Output:** U and V
- 3: *Initialization:* Model parameters U, V, Π and Σ are randomly initialized
- 4: In E-step posterior responsibility $\gamma_{ijk}^{(x)}$ is estimated using Eq. (8)
- 5: **For** Until convergence
- 6: (M-Step for updating $\Sigma^{(x+1)}$ and $\Pi^{(x+1)}$) Model parameters Σ and Π are computed using Eq. (9)
- 7: (M-Step for estimating V and U) Model parameters U and V are updated using Eq. (10)
- 8: (E-step for posterior responsibility γ_{ijk}) posterior responsibility γ_{ijk} is computed using current model parameters
- 9: **Return** User and Item latent factor matrices U and V

5 EXPERIMENTAL EVALUATION

In this section, we evaluate the effectiveness of our proposed recommendation model through real-world datasets.

5.1 Datasets

We use three datasets: Movielens [21], Libimseti [23] and Jester [22] in the evaluation. Table 2 provides a detailed view of the datasets. For privacy budget ϵ , we consider the value range from 0.1 to 3, lower values of privacy budget ϵ guarantee stronger privacy protection for users.

TABLE 2
Rating Datasets

Dataset	Total Ratings	No of Items	No of Users	Rating Scale
Movielens	100k	1682	943	0.5 to 5
Jester	2 Million	100	73,421	-10 to 10
Libimseti	17,359,346	168,791	135,359	1 to 10

5.2 Evaluation Metrics

5.2.1 RMSE

We use the Root Mean Squared Error (RMSE) to evaluate the predictive accuracy. We calculate the RMSE values over 10-fold cross-validation. RMSE can be estimated as :

$$RMSE = \sqrt{\frac{\sum_{i=0}^{n-1} (r_i - \hat{r}_i)^2}{n}},$$

where r_i is the actual rating, $\hat{r}_i$ is the predicted rating and n is the total number of ratings in the aggregated dataset.

5.2.2 F-Score

We use F-score to evaluate the utility of our recommendation system. Table 3 provides a detailed visualisation of how good a recommendation model is at predicting recommendations. Positives represent recommended items and, negatives represent non-recommended items.

TABLE 3
Confusion Matrix

Predicted Recommendations	Actual Recommendations		
		Positives	Negatives
	Positives	True Positive	False Positive
	Negatives	False Negative	True Negative

Precision and recall can be computed as follow:

$$precision = \frac{\#true\ positives}{\#true\ positives + \#false\ positives},$$

$$recall = \frac{\#true\ positives}{\#true\ positives + \#false\ negatives}.$$

We calculate F-score over top-10 recommended items. F-score can be computed as :

$$F-Score = 2 * \frac{Recall * Precision}{Recall + Precision}.$$

5.3 Results

5.3.1 Noise Distribution Evaluation

We derive the BLP noise distribution theoretically in section 4.1.2. In this section, we show that the noise distribution of Laplace and BLP mechanisms are distinct. We generate 100,000 random noise samples using BLP and Laplace mechanisms for the Movielens dataset while positioning their privacy budget ϵ to 0.1 and 1. Fig. 2a and 2b display the probability of noise samples drawn by Laplace and Bounded Laplace mechanisms. From the probability density functions, we note that the noise distribution of the two mechanisms is distinct. We also plot the BLP noise distribution curve based on our derived noise distribution expressions given by Eq. (4) and (5). Fig. 2a and 2b show that the theoretical derivation of distribution follows the experimental distributions exactly.

5.3.2 Influence of BLP on predictive accuracy

In this experiment, we demonstrate that using BLP as an input perturbation mechanism does play a significant role in obtaining higher predictive accuracy. We measure the RMSE when either BLP or Laplace act as the input perturbation mechanism while using the same rating prediction model (MoG or SVD). Fig. 3a and 3b display the resulting RMSE metric values for Movielens and Jester datasets respectively. The BLP mechanism results in higher recommendation accuracy than the Laplace mechanism for both rating prediction models.

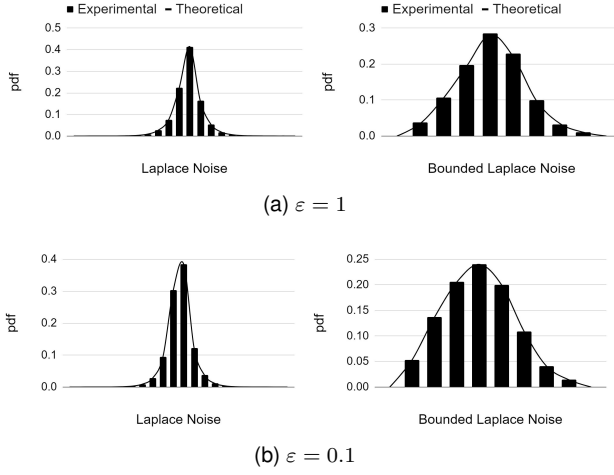


Fig. 2. Laplace vs Bounded Laplace Noise Distribution

5.3.3 Influence of MoG on predictive accuracy

In this experiment, we demonstrate that employing MoG in our recommendation model aids to improve predictive accuracy for lower values of privacy budget ϵ . We measure the RMSE values when using the MF with MoG or the SVD for rating prediction while using the same data perturbation mechanism. Fig. 4a and 4b display the resulting RMSE values for Movielens and Jester datasets respectively. For both datasets, the predictive accuracy from the MoG prediction model is much higher than SVD.

5.3.4 Predictive accuracy comparison over other private recommendation models

We compare the predictive accuracy of our recommendation model with other existing local differentially private recommendation models such as:

- Input Perturbation Method (ISGD) [6]: This method perturbs the user's original ratings locally using the Laplace mechanism. However, they apply a truncation method to ensure that the perturbed rating falls within a pre-defined domain. The noised ratings that fall out of a pre-defined range are clamped to either lower or upper bound of the rating domain using a threshold value. ISGD method uses MF for rating prediction at the SP side.
- Private Gradient-Matrix Factorization (PG-MF) [4]: This approach uses MF to perform recommendations. In this approach, the user computes user latent factors locally without submitting them to the SP. The SP estimates the item latent factors after collecting gradients from the users. Users, on the other hand, compute a perturbed gradient and submit that to the SP. The SP aggregates the perturbed gradient from all the users and then update the item latent factors accordingly. This method requires iterative communication between users and SP.

We use Non-Private MF as the baseline method as it does not use any local perturbation mechanism to perturb the user's original ratings. Instead, the MF algorithm uses actual ratings to predict missing ratings. The baseline method provides us with a lower bound RMSE value for predictive

TABLE 4
Comparison of Communication Cost for Movielens

Recommendation Model	User to SP	SP to User
BLP-MoG-MF	1 bit	No transfer
ISGD	1 bit	No transfer
PG-MF	1 bit	0.15MB

error. Our recommendation model (BLP-MoG-MF) uses BLP as input perturbation mechanism and MoG-MF as recommendation algorithm. The BLP-MoG-MF method uses the objective function specified by Eq. (1) to obtain latent factor matrices. ISGD and PG-MF methods also use the same objective function to perform rating predictions. To maintain the fairness of comparison, we did not compare our results with recommendation models that use different approaches to predict missing ratings.

Firstly, we compare BLP-MoG-MF with PG-MF. We vary the privacy budget ϵ from 0.1 to 1.6 for the Movielens dataset. Fig. 5 displays the RMSE values for BLP-MoG-MF, PG-MF and the baseline method. As expected, when the privacy budget increases, predictive accuracy for all privacy protection methods increases. Because when privacy budget ϵ increases, the magnitude of privacy loss LDP mechanism permits increases, which causes a rise in the predictive accuracy. More importantly, we notice that BLP-MoG-MF provides a lower RMSE than PG-MF for the same privacy budget ϵ .

Then, we compare BLP-MoG-MF with the ISGD method for Movielens, Libimseti and Jester datasets. We vary the privacy budget ϵ from 0.1 to 3 for all the datasets in this simulation. Fig. 7a, 7b and 7c display the RMSE values for BLP-MoG-MF, ISGD and the baseline methods. The results show that BLP-MoG-MF outperforms ISGD significantly for all values of the privacy budget ϵ . This trend implies that our method guarantees higher data utility for all the values of privacy budget ϵ .

5.3.5 Analysis of Communications Cost

We compare the communication cost incurred in our approach to recommendation models proposed by [4] and [6]. Table 4 summarises the analysis. Both BLP-MoG-MF and ISGD methods require the user to transfer a perturbed rating whenever the user rates an item. In the PG-MF method, the user transmits the perturbed gradient of user-latent factors to SP over multiple data transmission iterations. Both BLP-MoG-MF and ISGD methods do not require the SP to transmit any data back to the user. However, in the PG-MF approach at each iteration, the SP transmits an updated item latent factor matrix back to the user. This exchange between the SP and the user continues until the number of iterations reaches a pre-defined threshold value. We assume a single rating is 1 bit. The estimated size of the transmitted data for each iteration for the PG-MF method is approximately 0.15 MB for the Movielens dataset [1]. The comparison shows that we significantly reduce the communication cost in our proposed model compared to other local differential private recommendation models.

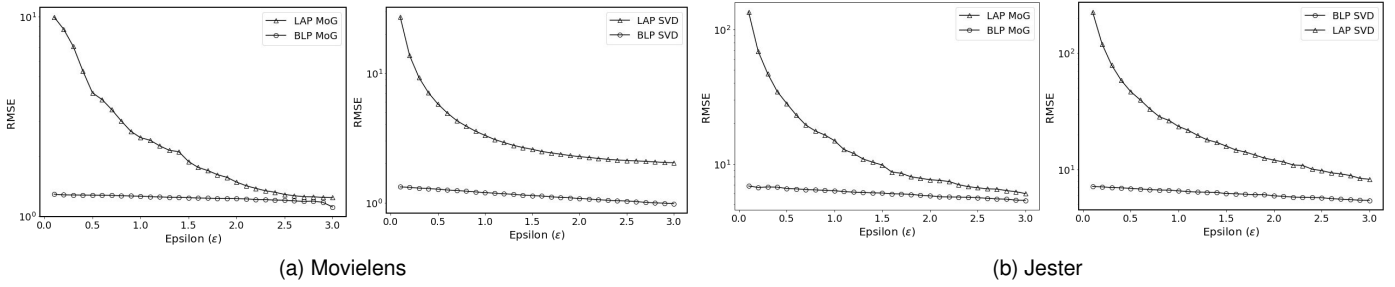


Fig. 3. Bounded Laplace Mechanism vs Laplace Mechanism RMSE Comparison

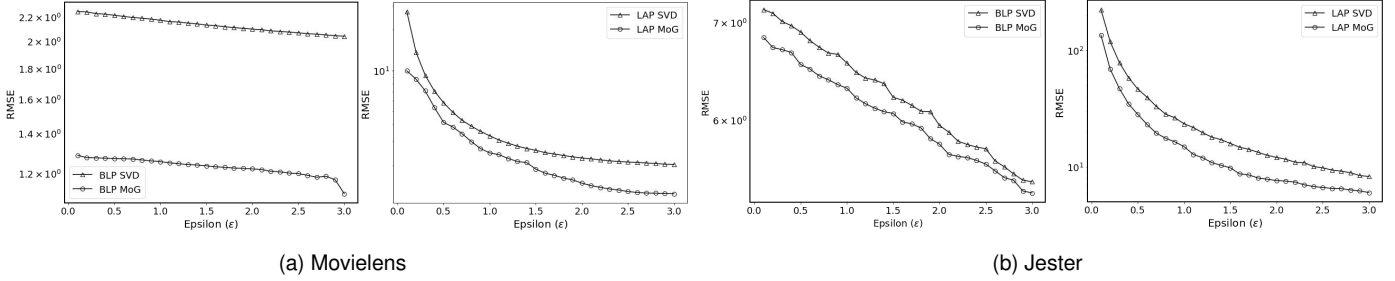


Fig. 4. MoG vs SVD Prediction Model RMSE Comparison

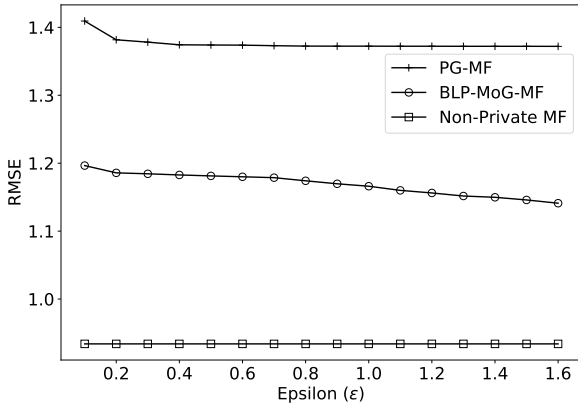


Fig. 5. PG-MF vs BLP-MoG-MF RMSE Comparison for MovieLens

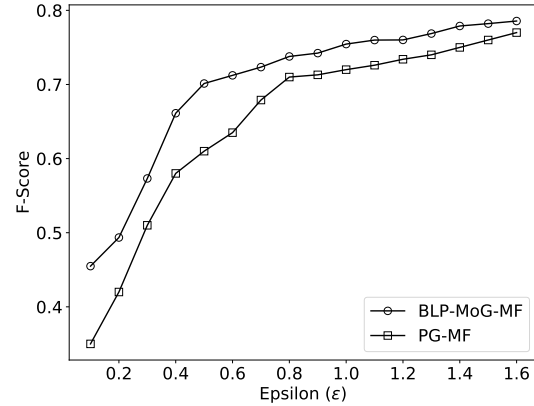


Fig. 6. PG-MF vs BLP-MoG-MF F-Score Comparison for MovieLens

5.3.6 Privacy Analysis

We use F-score as a utility metric to evaluate the recommendation accuracy and the privacy budget ϵ to measure the privacy loss of our system. We compute the F-score by comparing the top 10 items that our LDP based recommendation system recommends against the top 10 items recommended by other algorithms. Fig.6 demonstrates the F-score values for BLP-MoG-MF and PG-MF methods. Similar to RMSE values, the F-score value increases as the privacy budget ϵ increases. This increase indicates that the items recommended by the BLP-MoG-MF are becoming more similar to items recommended by the non-private recommendation system when the privacy loss is getting higher. The F-score results also show that BLP-MoG-MF provides more accurate recommendations compared to PG-MF for all values of privacy budget ϵ . Likewise, Fig. 8(a), 8(b) and 8(c) illustrate the F-score values for MovieLens, Jester and LibimSeti datasets

for BLP-MoG-MF and ISGD methods. There is a substantial increase in F-score as the privacy budget ϵ increases for all three datasets and both methods. Again, for all three datasets, the F-score of the BLP-MoG-MF is higher than the ISGD.

6 CONCLUSION

In our work, we have proposed a recommendation model under the consideration of an untrustworthy service provider. We have used BLP as a local input perturbation mechanism and MoG-MF for noise estimation and rating prediction. Compared to existing solutions, our proposed recommendation model can improve predictive accuracy and guarantees strong user privacy. Besides, our method does not incur any further communication cost to the user side as it only requires the user to transmit the perturbed rating to the SP.

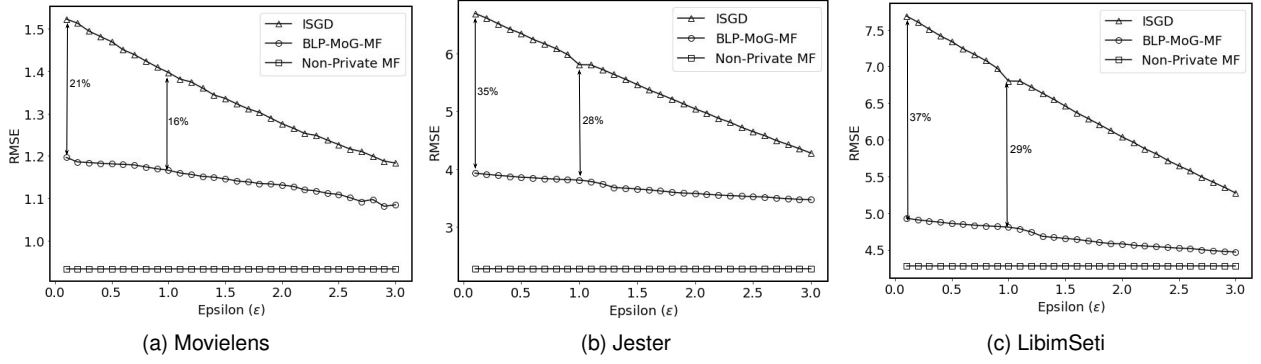


Fig. 7. BLP-MoG-MF vs ISGD RMSE Comparison

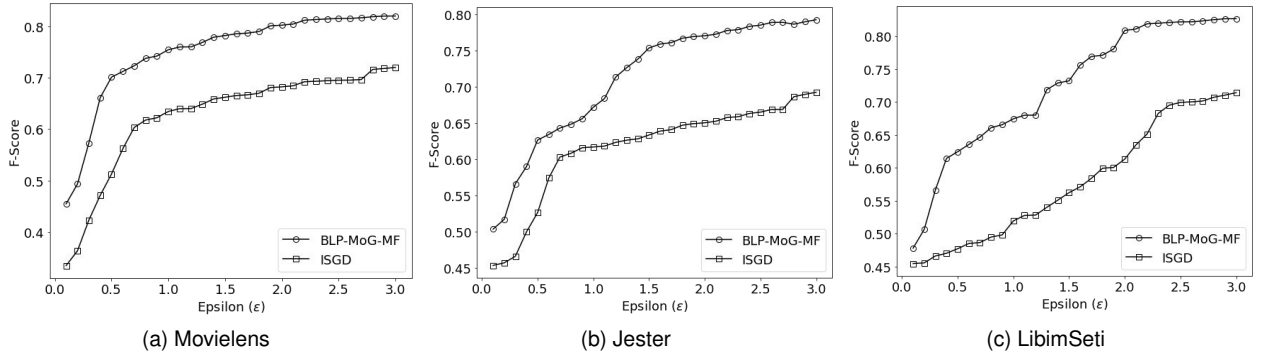


Fig. 8. BLP-MoG-MF vs ISGD RMSE Comparison

APPENDIX A

PROOF OF LEMMA 2

Assume r and r' are a pair of possible inputs of the BLP mechanism where $r' \geq r$ and $r' = r + z$. Let $0 \leq z \leq \Delta f$. In order to prove Lemma 2, we must first consider few other properties concerning $C(r)$. First we find $\frac{\partial}{\partial z} F(r, z) \geq 0$ when $r + z \leq u$.

$$\begin{aligned} \frac{\partial}{\partial z} F(r, z) &= \frac{1}{C(r)} \frac{\partial}{\partial z} \left(C(r+z) e^{\frac{z}{b}} \right) \\ &= \frac{1}{C(r)} \frac{\partial}{\partial z} \left(e^{\frac{z}{b}} - \frac{1}{2} \left(e^{-\frac{(r+z)-l}{b}} - e^{-\frac{u-(r+z)}{b}} \right) e^{\frac{z}{b}} \right) \\ &= \frac{1}{C(r)} \frac{\partial}{\partial z} \left(e^{\frac{z}{b}} - \frac{1}{2} \left(e^{-\frac{r+l}{b}} - e^{-\frac{u+r+2z}{b}} \right) \right) \\ &= \frac{1}{C(r)b} \left(1 - e^{-\frac{u-r-z}{b}} \right) e^{\frac{z}{b}} \end{aligned}$$

As $b > 0$, we then see that $\frac{\partial}{\partial z} F(r, z) \geq 0$ when $r + z \leq u$. Then we prove that $\frac{\partial}{\partial r} F(r, z) \leq 0$ when $z \geq 0$. First we note,

$$\frac{\partial}{\partial r} C(r+z) = \frac{1}{2b} \left(e^{-\frac{r+z-l}{b}} - e^{-\frac{u-r-z}{b}} \right)$$

We find,

$$\begin{aligned} \frac{\partial}{\partial r} F(r, z) &= \frac{e^{\frac{z}{b}}}{C(r)^2} \left(C(r) \frac{\partial}{\partial r} C(r+z) - C(r+z) \frac{\partial}{\partial r} C(r) \right) \\ &= \frac{e^{\frac{z}{b}}}{2bC(r)^2} \left(e^{-\frac{r-l}{b}} (e^{-\frac{z}{b}} - 1) + e^{-\frac{u-l-z}{b}} + e^{-\frac{u-r}{b}} (1 - e^{\frac{z}{b}}) \right. \\ &\quad \left. - e^{-\frac{u-l+z}{b}} \right) \end{aligned}$$

$$= \frac{e^{\frac{z}{b}} \left(\left(e^{-\frac{z}{b}} - 1 \right) \left(e^{\frac{u-r}{b}} - 1 \right) + \left(1 - e^{\frac{z}{b}} \right) \left(e^{\frac{r-l}{b}} - 1 \right) \right)}{2be^{\frac{u-l}{b}} C(r)^2}$$

Since $r \in [l, u]$, it proves that $e^{\frac{u-r}{b}}, e^{\frac{r-l}{b}} > 1$. When $z \geq 0$, it shows that $e^{-\frac{z}{b}} < 1$ and $e^{\frac{z}{b}} > 1$. Therefore, $\frac{\partial}{\partial r} F(r, z) \leq 0$ when $z \geq 0$.

As $\frac{\partial}{\partial r} F(r, z) \leq 0$, the maximum value of $F(r, z)$ at a fixed z^0 is attained at the smallest possible value of r , i.e $r = l$.

$$\max_{\substack{r, r+z^0 \in [l, u] \\ 0 \leq z^0 \leq \Delta f}} F(r, z^0) = \max_{0 \leq z^0 \leq \Delta f} \frac{C(l+z^0)}{C(l)} e^{\frac{z^0}{b}}$$

Then, as $\frac{\partial}{\partial z} F(l, z) \geq 0$, the maximum value of $F(l, z)$ is attained at the largest possible z , i.e $z = \Delta f$,

$$\max_{0 \leq z \leq \Delta f} \frac{C(l+z)}{C(l)} e^{\frac{z}{b}} = \frac{C(l+\Delta f)}{C(l)} e^{\frac{\Delta f}{b}}$$

REFERENCES

- [1] A. Narayanan and V. Shmatikov, "Robust De-anonymization of Large Sparse Datasets," 2008 IEEE Symposium on Security and Privacy, 2008.
- [2] F. Mcsherry and I. Mironov, "Differentially private recommender systems," Proceedings of the 15th ACM SIGKDD international conference on Knowledge discovery and data mining - KDD 09, 2009.

- [3] T. Zhu, G. Li, Y. Ren, W. Zhou, and P. Xiong, "Differential privacy for neighborhood-based collaborative filtering," Proceedings of the 2013 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining - ASONAM 13, 2013.
- [4] H. Shin, S. Kim, J. Shin, and X. Xiao, "Privacy Enhanced Matrix Factorization for Recommendation with Local Differential Privacy," IEEE Transactions on Knowledge and Data Engineering, vol. 30, no. 9, pp. 1770–1782, Jan. 2018.
- [5] J. Hua, C. Xia, and S. Zhong, "Differentially private matrix factorization," in Proc. 7th Int. Joint Conf. Artif. Intell., 2015, pp. 1763–1770.
- [6] A. Berlioz, A. Friedman, M. A. Kaafar, R. Boreli, and S. Berkovsky, "Applying Differential Privacy to Matrix Factorization," Proceedings of the 9th ACM Conference on Recommender Systems - RecSys 15, 2015.
- [7] Ú. Erlingsson, V. Pihur, and A. Korolova, "Rappor," Proceedings of the 2014 ACM SIGSAC Conference on Computer and Communications Security - CCS 14, 2014.
- [8] Apple's 'Differential Privacy' Is About Collecting Your Data—But Not Your Data (2016). [Online]. Available: <https://www.wired.com/2016/06/apples-differential-privacy-collecting-data/> (Accessed 24th Jan 2020).
- [9] X. Meng, "Personalized privacy-preserving social recommendation," in Proc. AAAI Conf. Artif. Intell., 2018.
- [10] S. P. Kasiviswanathan, H. K. Lee, K. Nissim, S. Raskhodnikova, and A. Smith, "What Can We Learn Privately?," 2008 49th Annual IEEE Symposium on Foundations of Computer Science, 2008.
- [11] F. Liu, "Generalized Gaussian Mechanism for Differential Privacy," IEEE Transactions on Knowledge and Data Engineering, vol. 31, no. 4, pp. 747–756, Jan. 2019.
- [12] R. Gemulla, E. Nijkamp, P. J. Haas, and Y. Sismanis, "Large-scale matrix factorization with distributed stochastic gradient descent," Proceedings of the 17th ACM SIGKDD international conference on Knowledge discovery and data mining - KDD 11, 2011.
- [13] G. Takács and D. Tikk, "Alternating least squares for personalized ranking," Proceedings of the sixth ACM conference on Recommender systems - RecSys 12, 2012.
- [14] M. Turk and A. Pentland, "Face recognition using eigenfaces," Proceedings. 1991 IEEE Computer Society Conference on Computer Vision and Pattern Recognition.
- [15] D. Meng and F. D. L. Torre, "Robust Matrix Factorization with Unknown Noise," 2013 IEEE International Conference on Computer Vision, 2013.
- [16] R. Vidal, R. Tron, and R. Hartley, "Multiframe Motion Segmentation with Missing Data Using Power Factorization and GPCA," International Journal of Computer Vision, vol. 79, no. 1, pp. 85–105, 2007.
- [17] A. P. Dempster, N. M. Laird, and D. B. Rubin, "Maximum likelihood from incomplete data via the EM algorithm," 1. Roy. Stat. Soc., vol. 39, no. 1, pp. 1–38, 1977.
- [18] F. de la Torre and M.J. Black, "A Framework for Robust Subspace Learning," Int'l J. Computer Vision, vol. 54, nos. 1–3, pp. 117–142, 2003.
- [19] A. Buchanan and A. Fitzgibbon, "Damped Newton Algorithms for Matrix Factorization with Missing Data," 2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR05).
- [20] N. Srebro and T. Jaakkola, "Weighted low-rank approximations," in Proc. ICML, 2003, pp. 720–727.
- [21] F. M. Harper and J. A. Konstan, "The MovieLens datasets: History and context," ACM Trans. Interact. Intell. Syst., vol. 5, no. 4, pp. 1–19, Dec. 2015.
- [22] K. Goldberg, T. Roeder, D. Gupta, and C. Perkins, "Eigentaste: A Constant Time Collaborative Filtering Algorithm," Information Retrieval J., vol. 4, no. 2, pp. 133–151, July 2001.
- [23] L. Brozovsky and V. Petricek, "Recommender system for online dating service," in Proc. 6th Conf. Zolosti, 2007, pp. 1–12.
- [24] K. Goldberg, T. Roeder, D. Gupta, and C. Perkins, "Eigentaste: A Constant Time Collaborative Filtering Algorithm," Information Retrieval J., vol. 4, no. 2, pp. 133–151, July 2001.
- [25] Y. Wang, X. Wu, and D. Hu, "Using randomized response for differential privacy preserving data collection," in Proc. EDBT/ICDT Workshops, 2016.
- [26] N. Holohan, S. Antonatos, S. Braghin and P. Mac Aonghusa, "The Bounded Laplace Mechanism in Differential Privacy," Journal of Privacy and Confidentiality, vol. 10, no. 1, 2019.
- [27] C. Cadwalladr and E. Graham-Harrison. (Mar. 2018). Revealed: 50 million Facebook profiles harvested for Cambridge analytica in major data breach. The Guardian. [Online]. Available: <https://www.theguardian.com/news/2018/mar/17/cambridge-analytica-facebook-influence-us-election> (Accessed: 25th April 2020).
- [28] R. Parameswaran and D. M. Blough, "Privacy preserving data obfuscation for inherently clustered data," Int. J. Inf. Comput. Secur., vol. 2, no. 1, pp. 4–26, 2008, doi: 10.1504/ijics.016819.
- [29] Y. K. Jain and S. K. Bhandare, "Min max normalization based data perturbation method for privacy protection," International Journal of Computer and Communication Technology, vol. 2, no. 8, pp. 45–50, 2011.
- [30] C. Dwork and A. Roth, "The algorithmic foundations of differential privacy," Found. Trends Theor. Comput. Sci., vol. 9, nos. 3–4, pp. 211–407, 2014.
- [31] Z. Liu, Y.-X. Wang, and A. J. Smola, "Fast differentially private matrix factorization," in Proc. ACM Conf. Recommender Syst., 2015, pp. 171–178.
- [32] P. Kairouz, S. Oh, and P. Viswanath, "Extremal mechanisms for local differential privacy," CoRR, vol. abs/1407.1338, 2014.
- [33] Y. Shen and H. Jin, "Privacy-preserving personalized recommendation: An instance-based approach via differential privacy," in Proc. IEEE Int. Conf. Data Mining, 2014, pp. 540–549.
- [34] I. Yakut and H. Polat, "PRIVACY-PRESERVING SVD-BASED COLLABORATIVE FILTERING ON PARTITIONED DATA", International Journal of Information Technology and Decision Making, vol. 09, no. 03, pp. 473–502, 2010. Available: 10.1142/s0219622010003919.
- [35] R. Bassily, U. Stemmer, A. G. Thakurta, *et al.* Practical locally private heavy hitters. In Advances in Neural Information Processing Systems, pages 2285–2293, 2017.
- [36] Z. Qin, Y. Yang, T. Yu, I. Khalil, X. Xiao, and K. Ren. Heavy hitter estimation over set-valued data with local differential privacy. In CCS, 2016.
- [37] T. Wang, J. Blocki, N. Li, and S. Jha. Locally differentially private protocols for frequency estimation. In USENIX Security Symposium, pages 729–745, 2017.
- [38] Z. Zhang, T. Wang, N. Li, S. He, and J. Chen, "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. ACM, 2018, pp. 212–229.
- [39] G. Lin, F. Liang, W. Pan, and Z. Ming, "Fedrec: Federated recommendation with explicit feedback," IEEE Intelligent Systems, pp. 1–1, 2020, doi:10.1109/MIS.2020.3017205.
- [40] D. Chai, L. Wang, K. Chen, and Q. Yang, "Secure federated matrix factorization," IEEE Intelligent Systems, pp. 1–1, 2020, doi:10.1109/MIS.2020.3014880.



Jeyamohan Neera is a PhD candidate in the Department of Computer and Information Sciences at Northumbria University since 2018. She received her BSc (2008) in Computer Engineering from University of Peradeniya, Sri Lanka and the MSc (2012) in Information Security from University of Salford, UK. Her current research is concerned with privacy in e-commerce systems. Her research interests are privacy protection, privacy enhanced AI and Cyber-security.



Xiaomin Chen is a Senior Lecturer in Department of Computer and Information Sciences, Northumbria University, UK. Prior to this she worked as a Research Associate in the Department of Computer Science, Loughborough University, UK during 09/2014 – 08/2015. From 12/2012 to 08/2013, she worked as a Post-doc Research Fellow in the Hamilton Institute, Maynooth University, Ireland, where she was awarded her PhD degree in Mathematics in Sep 2013. Her research is mainly focused on data privacy, cyber security, machine learning, wireless networking, resource allocation and optimisation.



Nauman Aslam is Professor in the Department of Computer and Information Science, Northumbria University, UK. Before joining Northumbria University as a Senior Lecturer in 2011, he worked as an Assistant Professor at Dalhousie University, Canada. He received his PhD in Engineering Mathematics from Dalhousie University, Canada in 2008. He is leading the Cyber Security and Network Systems research group at Northumbria University. His research interests cover diverse but interconnected areas related to communication networks. His current research efforts are focused at addressing problems related to wireless body area networks and IoT, network security, QoS-aware communication in industrial wireless sensor networks and application of Artificial Intelligence (AI) in communication networks. He has published over 100 papers in peer-reviewed journals and conference. Dr Nauman is a member of IEEE.



Kezhi Wang received his B.E. and M.E. degrees in School of Automation from Chongqing University, China, in 2008 and 2011, respectively. He received his Ph.D. degree in Engineering from the University of Warwick, U.K. in 2015. He was a Senior Research Officer in University of Essex, U.K. Currently he is a Senior Lecturer with Department of Computer and Information Sciences at Northumbria University, U.K. His research interests include wireless communications and machine learning.



Zhan Shu received his B.Eng. degree in Automation from Huazhong University of Science and Technology in 2003, and the Ph.D. degree in Control Engineering from The University of Hong Kong in 2008. He was a Postdoctoral Fellow in the Hamilton Institute, National University of Ireland from 2009 to 2011, Maynooth, and a Lecturer in Faculty of Engineering and Physical Sciences, University of Southampton from 2011 to 2019. Now, he is an Associate Professor in the Department of Electrical and Computer

Engineering, University of Alberta. He is a Senior Member of IEEE, Member of IET, and an invited reviewer of Mathematical Review of the American Mathematical Society. He serves as an Associate Editor for Mathematical Problems in Engineering, Asian Journal of Control, Journal of The Franklin Institute, Proc. IMechE, Part I: Journal of Systems and Control Engineering, IET Electronics Letters, IET Control Theory and Applications, IEEE Trans. Automatic Control, and a member of the IEEE Control Systems Society Conference Editorial Board. His current research interests include hybrid systems, positive systems, robust control, estimation and filtering, reinforcement learning, control applications in power electronics and systems biology.