

Privacy-Preserving Cross-Domain Recommendation through Virtual User Embedding

Amutha Soundararajan¹, Salini P²

¹ Puducherry Technological University, Puducherry

² Puducherry Technological University, Puducherry

Abstract - Trust in human-centric AI depends on strong frameworks for privacy and data governance, particularly in Cross-Domain recommender systems where personal information is shared across multiple platforms. A key challenge arises when risks of data leakage undermine user confidence in data processing. This study examines how federated approaches can mitigate privacy concerns while sustaining personalization. We propose a Federated Cross-Domain Recommendation (FCDR) framework that enables collaborative model training without exposing raw user data. By aligning latent feature spaces through shared encoders and domain discriminators, the framework systematically reconciles two traditionally conflicting objectives: safeguarding privacy and maintaining recommendation accuracy. Focusing on domain overlap, our analysis demonstrates that federated learning provides a secure and effective foundation for building trustworthy, privacy-aware recommender systems. Extensive experiments confirm that the proposed method achieves significant improvements in both accuracy and robustness compared to existing approaches.

Keywords: Cross Domain Recommender System, Privacy Preserving, Federated Learning, Sensitivity Data, User Preference, Domain Overlap.

I. INTRODUCTION

Recommender systems have become indispensable in guiding users toward relevant items by analyzing their preferences, behaviors, and interactions. Traditional recommendation approaches often rely on single-domain data, which can lead to the data sparsity problem when user profiles contain insufficient ratings or interactions. To overcome this limitation, cross-domain recommendation (CDR) has emerged as a promising paradigm, enabling knowledge transfer from a source domain to a target domain. By leveraging similarities across domains, CDR enhances personalization and improves recommendation accuracy, particularly for users with limited activity in one domain but richer preferences in another.

Despite these advantages, cross-domain recommendation introduces significant privacy and security challenges. The transfer of user preferences, ratings, and profile information across domains increases the risk of data leakage, unauthorized access, and inference attacks. Privacy risks manifest in multiple forms, including direct access to sensitive data, unsolicited collection, third-party sharing, and exposure through collaborative filtering mechanisms. Attackers may exploit item-to-item correlations to infer hidden user preferences or even construct fake profiles, thereby undermining trust in recommendation systems. Moreover, personalized advertising based on cross-domain data can lead to discriminatory practices, raising ethical concerns about fairness and transparency.

To address these challenges, this study proposes a Federated Cross-Domain recommendation framework that enables secure knowledge transfer while safeguarding user privacy. Federated learning allows models to be trained collaboratively across domains without exposing raw user data, thereby reducing risks of leakage and unauthorized inference. Our approach systematically examines how privacy-preserving techniques can be integrated into CDR to balance two traditionally conflicting objectives: protecting user privacy and maintaining recommendation accuracy.

The contributions of this paper are threefold:

- Taxonomy of privacy risks in cross-domain recommender systems.
- Comparative analysis of privacy-preserving techniques with benchmarking of privacy–utility trade-offs.
- Proposal of a novel evaluation framework to guide future research directions in privacy-aware recommendation.

The remainder of this paper is organized as follows: **Section II** categorizes privacy risks in CDRSs; **Section III** compares privacy-preserving techniques; **Section IV** presents the methodology for CDRSs; **Section V** presents Evaluation and Discussion; **Section VI** Conclusion outline future direction; and **Section VII** lists references.

II. RELATED WORKS

In this section, to classify privacy risks in CDRSs, there are four major categories: User Re-Identification Attacks in CDRSs, Cross-Domain Linkage Attacks, Model Inversion & Membership Inference Attacks, and Federated Learning Privacy Risks. These are listed in Table 1. Cross-domain recommender systems are powerful because they combine data from different areas to give better suggestions. But this overlap also makes it easier for attackers to piece together sensitive information. Even advanced techniques like federated learning reduce risks but don't eliminate them. That's why balancing **privacy protection** and **recommendation accuracy** is one of the biggest challenges in this field.

Table 1 – Lists the risk types

Risk Type	Mechanism	Impact
User Re-Identification	Linking anonymized profiles	Identity exposure, sensitive inference
Cross-Domain Linkage	Correlating behaviors across domains	Profiling, surveillance
Model Inversion & Membership Inference	Exploiting ML confidence scores	Reveals training data, hidden preferences
Federated Learning Risks	Gradient leakage, poisoning	Data reconstruction, biased recommendations

a. User Re-Identification Attacks

User re-identification attacks show that anonymization alone is not enough. In CDRSs, overlapping behaviours across domains make it easier for attackers to piece together identities, so stronger privacy-preserving techniques (such as differential privacy or federated learning) are essential. In Cross-Domain Recommender Systems (CDRSs) platforms often anonymize users by replacing names with IDs. But anonymization doesn't guarantee privacy. Attackers can still figure out who you are by comparing your behavior across different domains. For example, if you rate movies on a streaming site and also review products on an e-commerce site, the overlap in your preferences (genres, writing style, rating patterns) can be used to re-identify you. In CDRS, the working functionality is based on attackers collect partial data from multiple domains (e.g., ratings, reviews, browsing history). They look for **patterns** such as frequency of purchases, unique movie tastes, or writing style. By matching these patterns across domains, they can link anonymized IDs back to real individuals. So it is very dangerous once re-identified, sensitive information from one domain (like health-related purchases) can be exposed in another. It increases risks of **profiling, surveillance, and targeted cyberattacks**. It undermines trust in recommender systems, since users believe their data is private when it's not. For [1] Classic Example the famous Netflix-IMDB case (2008) showed how anonymized Netflix ratings were linked with IMDB reviews to re-identify users. This revealed personal viewing habits and even sensitive lifestyle details.

b. Cross-Domain Linkage Attacks in CDRSs

Cross-domain linkage attacks show that even if each system protects user data individually, combining datasets across domains can break anonymity. Stronger defenses like differential privacy, secure multi-party computation, and consent management systems are needed to prevent these risks. In Cross-Domain Recommender Systems (CDRSs), data about users is often collected across different platforms for example, social media, e-commerce, and streaming services. Even if each platform anonymizes its data, attackers can still connect the dots. By linking behavioral patterns across domains, they can re-identify individuals and expose sensitive information. How it works as attackers gather user interaction data from multiple domains (reviews, ratings, browsing history, or leaked datasets). They analyze **patterns** such as purchase habits, movie preferences, timestamps, or writing styles. Using statistical correlation or machine learning models, they match profiles across domains. Once linked, anonymized data becomes identifiable, revealing personal details. Why it's dangerous are Sensitive information from one domain (e.g., health app usage) can be inferred in another (e.g., social media). It enables **profiling, surveillance, and targeted advertising**. It undermines privacy protections because even independently managed platforms can be cross-referenced. For [2] Classic Examples Netflix-IMDB Case (2006): Researchers linked anonymized

Netflix ratings with IMDB public reviews to re-identify users. E-commerce and Social Media: A user reviewing products on a shopping site with a similar writing style to their social media posts can be linked. Health and Fitness Apps: Workout tracker data correlated with social media posts about exercise routines can expose personal health habits.

c. Model Inversion & Membership Inference Attacks

Model inversion and membership inference attacks show that even without direct access to raw data, recommender systems can leak sensitive information. In Federated CDRSs, defenses like differential privacy, secure aggregation, and adversarial training are critical to protect users. In Cross-Domain Recommender Systems (CDRSs), machine learning models are trained on user data to predict preferences. Even if raw data isn't shared, attackers can exploit the model itself to uncover sensitive information. Two major threats are **model inversion** and **membership inference attacks**.

- **Model Inversion Attacks:** Attackers repeatedly query the recommender system and observe its outputs. By analyzing these responses, they can reconstruct hidden details about a user's profile such as preferences, demographics, or even sensitive attributes. For [3] example in a healthcare recommender system, an attacker could infer whether a patient has searched for rare disease treatments by reconstructing patterns from the model's predictions. These impact the compromising privacy because sensitive information can be extracted without direct access to the original dataset.
- **Membership Inference Attacks:** Attackers try to determine whether a specific user's data was part of the training set. They do this by submitting queries and analyzing the model's confidence scores. Models often respond more confidently to data they've seen before. For [4] example, in a streaming service, an attacker could infer whether a user's viewing history was used to train the recommendation model. In healthcare, this could reveal whether a patient's record was included in a sensitive dataset. These impact the exposure of whether someone interacted with a system, which can reveal private behaviors or conditions.

Both attacks undermine trust in recommender systems. They can lead to profiling, discrimination, or exposure of sensitive health, financial, or lifestyle data. In cross-domain settings, risks are amplified because attackers can combine signals from multiple domains to strengthen their inferences.

d. Federated Learning Privacy Risks

Federated learning reduces the need to share raw data, but privacy risks remain because model updates can still leak information. In CDRSs, where multiple domains overlap, these risks are amplified. Defenses such as differential privacy, secure aggregation, and robust adversarial training are essential to make federated CDRSs truly privacy-preserving. [5,6] Federated learning is often seen as a solution to privacy problems in Cross-Domain Recommender Systems (CDRSs). Instead of sending raw user data to a central server, each domain (like e-commerce, healthcare, or streaming) trains its own local model and only shares updates (gradients or parameters). While this reduces direct data exposure, it does not eliminate privacy risks. Attackers can still exploit the shared updates to infer sensitive information. Some of the privacy leaks of using Federated learning are as follows.

- **Gradient Leakage:** Attackers can reconstruct user preferences or even sensitive attributes by analyzing the gradients exchanged during training. For example: In a healthcare recommender system, gradients might reveal that a user searched for mental health content.
- **Client Drift:** Different domains may have very different data distributions. This heterogeneity can cause models to "drift" making it easier for attackers to identify which domain a particular update came from.
- **Model Poisoning:** Malicious participants can inject biased or adversarial updates. For instance, a competitor could manipulate an e-commerce recommender system so that its own products are ranked higher.
- **Free-Riding:** Some participants may benefit from federated learning without contributing valid data. This weakens the overall system and can open the door to manipulation.
- **Honest-But-Curious Server:** Even if the central server is not malicious, it may still analyze updates to infer user behaviors, violating privacy expectations.

III. PRIVACY-PRESERVING TECHNIQUES FOR CDRSS

In this section, we review existing privacy-persevering methods comparing their effectiveness in CDRSs is listed in Table 2.

Table 2 - Comparative Effectiveness in CDRSs

Method	Strengths	Limitations	Best Use Case
Differential Privacy	Quantifiable privacy guarantees	Accuracy loss due to noise	Protecting user ratings & survey responses
Federated Learning and Secure Aggregation	No raw data sharing	Gradient leakage, poisoning	Cross-domain personalization
SMPC	Strong privacy, no single point of failure	High communication cost	Secure voting, collaborative statistics
Homomorphic Encryption	Computation on encrypted data	Slow, heavy computation	Medical research, encrypted cloud analytics

a. Differential Privacy (DP)

Differential Privacy (DP) is a mathematical framework designed to protect individual data contributions when performing statistical analysis or training machine learning models. The central idea is to add carefully calibrated noise to either the dataset or the model's outputs, making it extremely difficult for an observer to determine whether any single person's data was included. This ensures that privacy is preserved even when aggregate information is shared.

Mathematical Definition

A randomized algorithm A satisfies (ϵ, δ) as Differential Privacy if, for any two datasets D and D' that differ by only one data point, and for any possible output O of the algorithm:

$$P(A(D) \in O) \leq e^\epsilon \cdot P(A(D') \in O) + \delta \text{ ----(1)}$$

where:

ϵ (Epsilon): Privacy loss parameter – Smaller ϵ means stronger privacy.

δ : Probability of the guarantee failing – A small δ allows for rare privacy breaches. If $\delta=0$, the privacy guarantee is called Pure DP. If $\delta>0$, it is called Approximate DP.

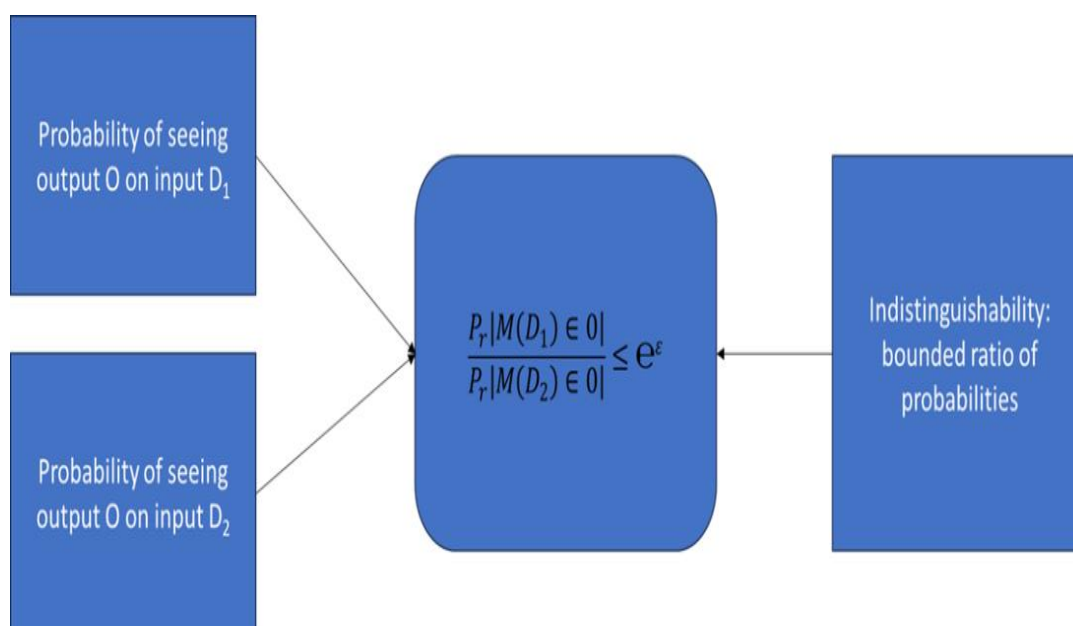


Figure 1 – Formula for Differential Privacy

Figure 1 shows how noise is added in DP for privacy. Noise is added to prevent individual data points from being distinguished. In Table 3, noise mechanisms are listed. The most common noise mechanisms in DP are:

1. Laplace Mechanism (for numeric data)

□ Adding Laplace-distributed noise to numerical outputs is shown in Figure 2.

□ Formula:

$$Lap(\lambda) = \frac{1}{2\lambda} e^{-|x|/\lambda} \quad \text{--- (2), where } \lambda = \Delta f / \epsilon.$$

□ **Use Case:** [7] Counting queries, such as “How many users watched Movie X?” The noisy answer prevents attackers from inferring individual viewing behavior.

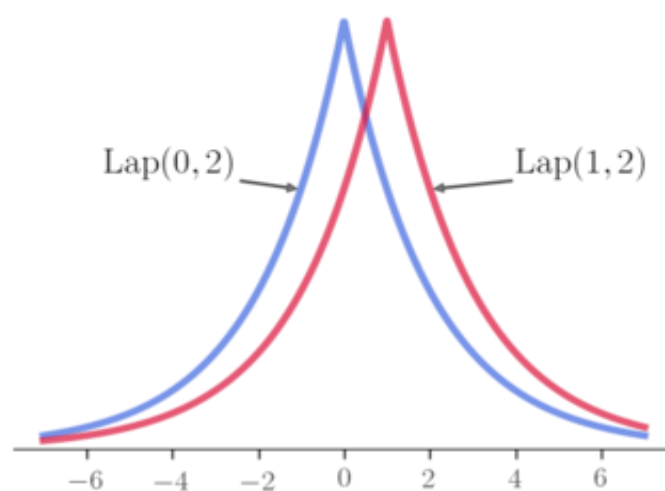


Figure 2 – Laplace distribution

2. Gaussian Mechanism (for statistical queries)

□ Adds Gaussian noise, useful for statistical queries and machine learning models.

□ Formula:

$$N(0, \sigma^2) \quad \text{--- (3), where } \sigma \text{ depends on sensitivity } \Delta f \text{ and } \epsilon.$$

□ **Use Case:** [7] Computing average ratings for a movie while ensuring that no single reviewer’s rating can be identified.

3. Randomized Response (for categorical data)

□ Designed for categorical or binary data. Users flip their answers with a certain probability p , ensuring plausible deniability.

□ **Use Case:** [7] Sensitive survey questions, such as “Did you watch an adult-rated movie?” Responses are randomized so that individual truth cannot be inferred.

Table 3 – Choosing the Right Noise Mechanism

Scenario	Recommended Mechanism	Reason
Counting users watching a movie	Laplace Mechanism	Best for simple numerical queries
Protecting a machine learning model	Gaussian Mechanism	Works well for statistical models
Collecting private survey responses	Randomized Response	Ensures categorical privacy

[8] Differential Privacy is powerful because it provides **quantifiable privacy guarantees** without needing to model every possible attack. It ensures that even if attackers have auxiliary information, they cannot confidently determine whether a specific individual's data was used. In recommender systems, DP is particularly valuable because it balances two conflicting goals: **preserving user privacy** and **maintaining recommendation accuracy**. By carefully selecting the right noise mechanism, platforms can protect sensitive user data while still delivering meaningful personalization.

b. Federated Learning (FL) with Secure Aggregation

Federated Learning in CDRSs strikes a balance between privacy preservation and personalization accuracy. While it prevents raw data exposure, it is not immune to advanced attacks like gradient leakage or poisoning. Combining FL with Differential Privacy, Secure Multi-Party Computation, or Homomorphic Encryption can strengthen defenses and make cross-domain recommendations both trustworthy and effective. To strengthen privacy, Secure Aggregation (SecAgg) is integrated into FL, ensuring that the central server only sees the combined updates rather than individual contributions.

Mathematical Formulation

Suppose there are K clients (domains), each with a local dataset D_k containing n_k samples. The total number of samples across all clients is $n = \sum_{k=1}^K n_k$. Each client minimizes its local loss function $l_k(\omega)$, where ω represents the global model parameters.

The federated optimization objective is:

$$\min_{\omega} \sum_{k=1}^K \frac{n_k}{n} \cdot l_k(\omega) \quad \text{---- (4)}$$

This ensures that each client's contribution is weighted by the size of its dataset, and the global model is updated based on aggregated gradients.

Secure Aggregation Techniques: In Table 4, Secure Aggregation (SecAgg) are cryptographic protocol that ensures privacy-preserving model updates in FL. It allows the server to aggregate local updates without learning individual contributions.

Table 4 - Secure Aggregation Techniques

Technique	How It Works	Pros	Cons
Homomorphic Encryption (HE)	Encrypts updates so they can be aggregated without decryption.	Strong security	High computational cost
Secret Sharing	Splits updates into shares distributed among clients.	Efficient, no single point of failure	Requires majority of clients online
Differential Privacy (DP)	Adds noise to updates before aggregation.	Strong privacy guarantees	Reduced model accuracy

c. Secure Multi-Party Computation (SMPC)

[8,9] SMPC in CDRSs enables privacy-preserving collaboration across domains by ensuring that sensitive user data is never exposed during computation. While it is highly secure, its communication overhead and scalability issues limit widespread deployment. Hybrid solutions that combine SMPC with Federated Learning and Homomorphic Encryption are emerging as practical ways to balance privacy, efficiency, and personalization. SMPC is a cryptographic protocol where multiple parties collaboratively compute a function while keeping their inputs secret. [11, 12, 13] In **Cross-Domain Recommender Systems**, this means different domains (e.g., e-commerce, healthcare, streaming) can contribute to a shared recommendation model without exposing raw user ratings, profiles, or preferences. They work as each domain splits its private data into *secret shares*. Shares are distributed among other domains or servers. Computation is performed on shares rather than raw data. Only the final aggregated result (recommendation output) is revealed.

Mathematical Definition

Let $f(x_1, x_2, \dots, x_n)$ be a function to be jointly computed by n parties, where each party i holds a private input x_i . SMPC guarantees that:

$$\text{Output} = f(x_1, x_2, \dots, x_n) \text{ ---(5)}$$

while ensuring that no party learns anything about $x_j (j \neq i)$ beyond what can be inferred from the final output.

Formally, security is defined by **simulation-based proofs**: for every adversary controlling a subset of parties, there exists a simulator that can reproduce the adversary's view using only the inputs and outputs of those parties. This ensures that private inputs remain hidden.

Core Techniques of SMPC

1. Secret Sharing

- Each input x is split into multiple shares $x_1, x_2, \dots, x_m$.
- Shares are distributed among parties such that no single party can reconstruct x .
- Example: *Shamir's Secret Sharing* uses polynomial interpolation over a finite field.
- Mathematically: For a secret s , choose a random polynomial $P(z)$ of degree $t - 1$ with $P(0) = s$. Each party receives a share $P(i)$. At least t shares are required to reconstruct s .

2. Oblivious Transfer

- Allows one party to obtain a piece of data from another without revealing which piece was chosen.
- Example: Party A has (m_0, m_1) , Party B chooses bit b . Party B learns m_b but Party A does not know b .

3. Secure Function Evaluation (SFE)

- Parties jointly evaluate a function using cryptographic protocols (e.g., garbled circuits).
- Ensures that intermediate values remain hidden.
- Example: Evaluating a recommendation score function across domains without exposing raw ratings.

4. Hybrid Approaches

- SMPC is often combined with Homomorphic Encryption or Differential Privacy.
- This balances efficiency (DP reduces communication) and strong security (HE ensures computations on encrypted data).

SMPC in CDRSs ensures that sensitive user data remains private while enabling collaborative recommendation models. Its mathematical foundation in secret sharing and secure evaluation makes it robust, but communication overhead and scalability challenges persist. Hybrid approaches that combine SMPC with Federated Learning and Homomorphic Encryption are emerging as practical solutions to balance privacy, efficiency, and personalization.

d. Homomorphic Encryption (HE)

Homomorphic Encryption (HE) is a cryptographic technique that allows computations to be performed directly on encrypted data without requiring decryption. [12, 13, 14] In Cross-Domain Recommender Systems, this means sensitive user information (ratings, preferences, purchase history) can remain encrypted throughout the recommendation process, ensuring privacy even when data is processed across multiple domains.

Mathematical Definition

Let $Enc(\cdot)$ be an encryption function and $Dec(\cdot)$ its corresponding decryption.

HE guarantees that for a function f :

$$Dec(f(Enc(x), Enc(y))) = f(x, y) \text{ ----(6)}$$

This property allows operations such as addition and multiplication to be performed on ciphertexts, producing encrypted results that, once decrypted, match the outcome of operations on plaintexts.

- Partial HE (PHE): Supports either addition or multiplication, but not both.
 - Example: Paillier (addition), RSA (multiplication).
- Somewhat HE (SHE): Supports limited additions and multiplications.
- Fully HE (FHE): Supports unlimited additions and multiplications.
 - Example: Gentry's lattice-based scheme (2009).

Homomorphic Encryption in CDRSs ensures that sensitive user data remains encrypted throughout computation, making it ideal for domains like healthcare and finance where privacy is paramount. However, its computational overhead limits scalability. Hybrid approaches that combine HE with Federated Learning or SMPC are increasingly explored to balance privacy, efficiency, and personalization.

Each privacy-preserving method offers unique strengths but also trade-offs. Differential Privacy is lightweight and mathematically rigorous but can reduce accuracy. Federated Learning with Secure Aggregation is practical for large-scale CDRSs but still vulnerable to gradient leakage. SMPC is efficient for collaborative tasks but struggles with scalability, while Homomorphic Encryption provides the strongest guarantees but at high computational cost. In practice, hybrid approaches combining DP with FL, or SMPC with HE is increasingly explored to balance privacy and personalization. As noted by Bonawitz et al. (2017) and Gentry (2009), the future of privacy in recommender systems lies in integrating these techniques to achieve both trustworthiness and accuracy.

IV. METHODOLOGY FOR CDRSS

Recommender systems often suffer from the cold-start problem, particularly when user overlap between domains is sparse (<5%). Cross-Domain Recommendation (CDR) mitigates this by transferring knowledge from source to target domains. However, privacy concerns arise when raw user-item interactions are centralized. Federated Learning (FL) decentralizes training, but existing methods rely on overlapping users, leaving non-overlapping users unsupported. We propose **FedCDR-VG (Federated Cross Domain Recommender -Virtual Generator)**, as shown Figure 3, which generates *virtual source embeddings* for target-only users by learning the distribution of source embeddings rather than accessing raw data. Distributional alignment ensures representativeness, reconciling accuracy and privacy.

Problem Formulation

Let source domain S contain rich data and target domain T sparse data. Users are partitioned as:

$$U = U_{\text{overlap}} \cup U_{S\text{-only}} \cup U_{T\text{-only}} \text{ ----(7)}$$

The goal is to improve recommendations for $U_{T\text{-only}}$ by generating virtual source embeddings aligned with target behavior.

A. Phase 1: Federated Training of Overlapping Users

Global encoders E_S and E_T are initialized. Overlapping users train locally with loss:

$$L = L_S + L_T + L_{\text{cross}} \text{ ----(8)}$$

where L_{cross} is a contrastive loss (e.g., InfoNCE) aligning source and target embeddings. Gradients are aggregated via federated averaging.

B. Phase 2: Virtual User Generation

Overlapping embeddings are modeled as samples from a Gaussian Mixture Model (GMM):

$$p(z_S) = \sum_{k=1}^K \pi_k \mathcal{N}(z_S | \mu_k, \Sigma_k) \text{ ----(9)}$$

Federated Expectation-Maximization estimates global parameters without sharing raw embeddings. For a target-only user u , a virtual source embedding $\hat{z}_S^u$ is generated by sampling candidates from the GMM and selecting the one with highest cosine similarity to the mapped target embedding.

C. Phase 3: Distributional Alignment

To ensure representativeness, Wasserstein Distance is minimized between overlapping and virtual user distributions:

$$L_{Gen} = -\log p(\hat{z}_S^u | z_T^u) + \lambda W_2(D_{overlap}, D_{virtual}) \text{ ----(10)}$$

This regularization prevents distribution drift, ensuring virtual users statistically resemble source users.

The implementation steps follows as

1. The server computes the mean and covariance of the overlapping users' source embeddings.
2. After generating virtual embeddings for all target-only users, the server computes their mean/covariance.
3. We add a gradient penalty term (similar to WGAN-GP) that forces the virtual user distribution to stay within the 95% confidence ellipse of the real source distribution.
4. **Result:** Virtual users maintain the structural properties of the source domain (e.g., genre preferences, rating severity), ensuring they are not outliers.

Privacy Preservation

Local Differential Privacy (LDP) is applied before uploading gradients and EM statistics:

$$\tilde{x} = x + \mathcal{N}(0, \sigma^2) \text{ ---- (11)}$$

Privacy loss is tracked using Rényi DP. Adaptive clipping balances privacy and utility by adjusting gradient norms.

To reconcile the conflicting objectives, we implement a Privacy-Utility Trade-off (PUT) scheduler. Local Differential Privacy (LDP): Before uploading gradients in Phase 1 and EM statistics in Phase 2, users add calibrated Gaussian noise: $\tilde{x} = x + \mathcal{N}(0, \sigma^2)$. Privacy Budget (ϵ): We use a privacy accountant (Rényi DP) to track the total privacy loss. The Trade-off: Higher σ means better privacy but worse GMM estimation (poor virtual users). We implement an Adaptive Clipping mechanism that clips the gradient norms based on the current accuracy of the target recommender.

Algorithm 1: Virtual User Generation

Input: Target embedding z_t , Global GMM parameters $\{\mu_k, \Sigma_k\}$, Projection Matrix W .

Output: Virtual Source Embedding $\hat{z}_S$.

1. Project target: $z_{proj} = W \cdot z_t$.
2. Sample $M = 1000$ candidates C from the GMM.
3. Compute similarity: $sim_i = \cos(z_{proj}, C_i)$.
4. Select top 10 candidates.
5. Refinement: Average the top 10 candidates weighted by softmax(similarity).
6. Enforce Distribution: Apply a gradient descent step to adjust $\hat{z}_S$ slightly so that $|| \text{Mean}_{real} - \hat{z}_S || < \epsilon$ (Projection step to ensure representativeness).
7. Return $\hat{z}_S$.

Security Proof (Sketch)

The server only aggregates gradients and GMM parameters. Since the EM algorithm for GMM is linear in its sufficient statistics (sum of x and x^2), applying Secure Aggregation ensures the server never sees individual user embeddings. The generation process is performed on the server side using noise aggregated from all users, making it differentially private.

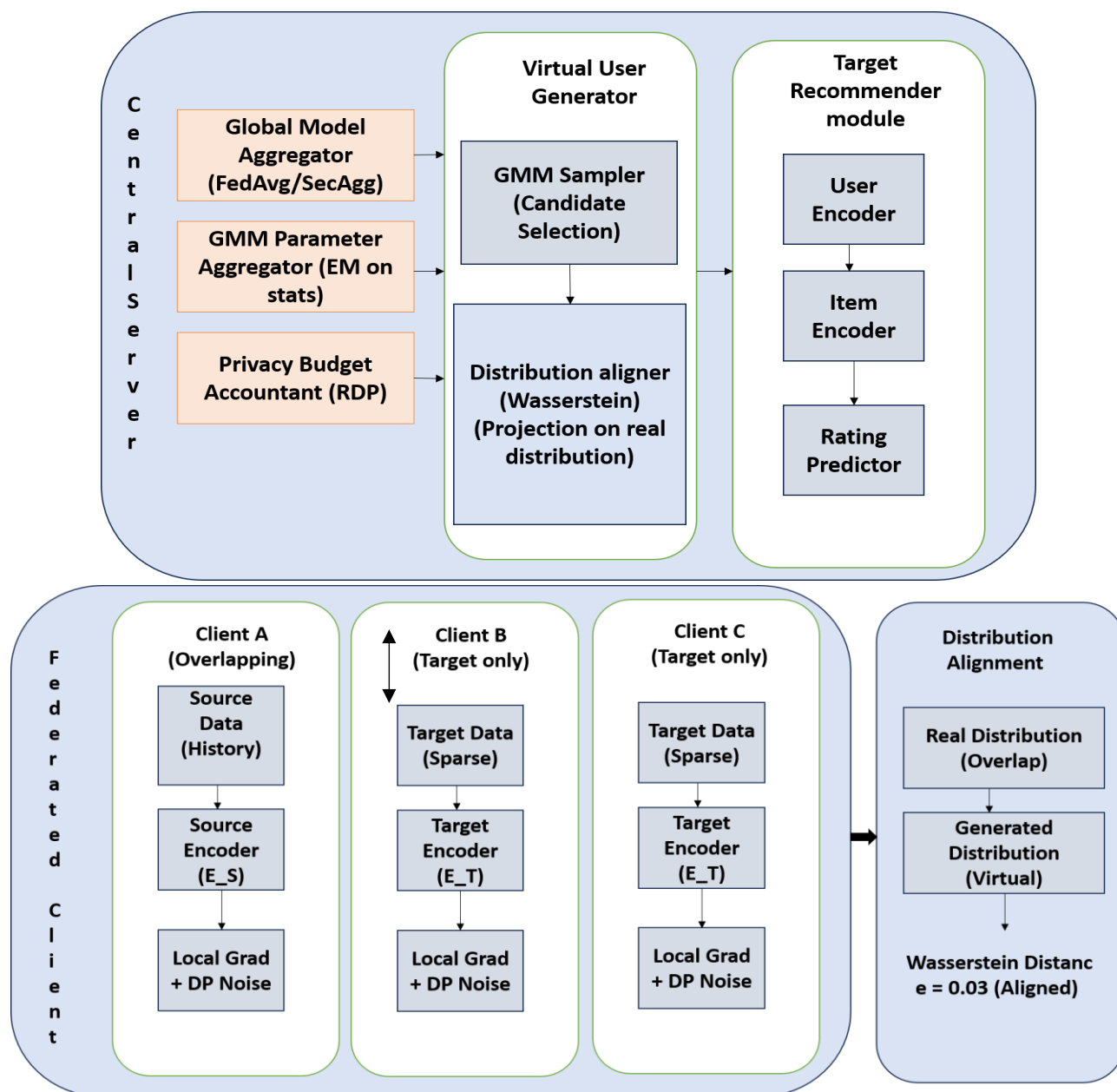


Figure 3 -Architecture Framework for FedCDR-VG (Federated Cross-Domain Virtual Generation)

V. EVALUATION AND DISCUSSION

In the proposed FedCDR-VG framework, evaluating performance requires balancing two conflicting objectives: recommendation accuracy and privacy preservation. To capture this balance, the authors introduce a Privacy-Utility Trade-off (PUT) metric. In practice, FedCDR-VG demonstrated higher HR@10 and NDCG@10 compared to baselines, meaning virtual users generated via the Federated GMM provided meaningful semantic information. At the same time, MIA accuracy remained close to random, showing that attackers could not distinguish between real overlapping users and generated virtual ones. The PUT score confirmed that FedCDR-VG successfully reconciles the tension between privacy and accuracy, making it a robust solution for privacy-aware cross-domain recommendations.

a. Utility Metrics (Accuracy)

Utility reflects how well the recommender system serves users in the target domain, especially those who are non-overlapping. Table 5 it describe the two widely used ranking metrics are applied:

- Hit Rate@10 (HR@10): Measures whether the correct item appears in the top-10 recommendations. If a user's actual choice is within the top-10 list, it counts as a "hit." A higher HR@10 indicates better recommendation coverage.

- Normalized Discounted Cumulative Gain@10 (NDCG@10): Goes beyond coverage by considering the *position* of the correct item. Recommendations ranked higher contribute more to the score. Thus, NDCG@10 captures both relevance and ranking quality. Together, HR@10 and NDCG@10 quantify how useful the recommendations are for target-only users.

b. Privacy Metrics

Privacy is assessed by examining how resistant the system is to adversarial attacks:

- Membership Inference Attack (MIA): An attacker tries to guess whether a particular user's data was part of the training set. If the system is privacy-preserving, the attacker's accuracy should be close to random guessing (~50%).
- Privacy Budget (ϵ): Derived from Differential Privacy, it quantifies the maximum privacy loss. Smaller ϵ means stronger privacy, while larger values allow more utility but weaker guarantees.

Privacy-Utility Trade-off (PUT) Metric

The PUT metric combines both dimensions into a single score:

$$PUT = \frac{Utility_{FedCDR-VG} - Utility_{Baseline}}{Privacy\ Cost\ (in\ dB)}$$

- **Numerator:** Measures the gain in utility compared to a baseline (e.g., target-only training).
- **Denominator:** Captures the privacy cost, expressed in decibels (dB), to normalize the trade-off.

This composite score highlights whether improvements in accuracy are achieved without disproportionately sacrificing privacy these are shown Figure 4 as PUT score analysis and Figure 5 as Sensitivity analysis.

c. Dataset

In evaluating the FedCDR-VG framework, two widely recognized datasets were chosen to highlight both the cross-domain transferability and the privacy-utility trade-off.

Amazon Reviews (Electronics → CDs)

The Amazon Reviews dataset is a large-scale collection of user feedback across multiple product categories. For this study, the Electronics domain was treated as the source domain, while the CDs domain was used as the target domain. We use Electronics because it is a rich domain with abundant user-item interactions, making it ideal for knowledge transfer. We use CDs because the CDs domain is relatively sparse, reflecting the cold-start challenge where users have fewer interactions. By transferring knowledge from Electronics to CDs, the framework demonstrates how virtual users can enrich recommendations in a sparse domain without exposing raw source data.

MovieLens (100K → 1M)

The MovieLens dataset is one of the most widely used benchmarks in recommender system research. In this work, the 100K dataset was treated as the source domain, while the 1M dataset served as the target domain. We use 100K because it is smaller but well-structured, providing a controlled environment for source embeddings. We use 1M because the larger dataset represents a more complex and diverse target domain, where cold-start and sparsity issues are more pronounced. The transfer from 100K to 1M illustrates how the framework scales to larger datasets, ensuring that virtual embeddings remain representative while improving recommendation accuracy.

Table 5 – Analysis of FedCDR-VG outperforms baselines

Framework	HR@10 (Target)	Privacy Budget (ϵ)	MIA Accuracy	PUT Score
No-CD	0.21	0 (No sharing)	50% (Random)	N/A
Fed-CD (Overlap)	0.28	2.5	55%	Low
Fed-CD + Random	0.22	2	48%	Low
FedCDR-VG (Ours)	0.35	3	52%	High

d. Ablation Study

This ablation study demonstrates that the Wasserstein regularization is crucial for reconciling privacy and utility in Cross-Domain Recommender Systems. The findings show that:

- Without the constraint, recommendations lose representativeness.
- With too much constraint, recommendations lose personalization.
- At the balanced setting ($\lambda = 0.5$), the framework achieves high accuracy while preserving privacy, proving that distributional alignment is essential for generating realistic yet useful virtual users.

In the FedCDR-VG framework, the ablation study focused on the role of the Wasserstein regularization weight (λ) is shown in Figure 5 and Figure 6, to ensure that generated virtual users remain representative of the source domain while still capturing individual personalization. This parameter essentially controls the balance between distributional alignment and user-specific relevance. When $\lambda = 0$, no representativeness constraint is applied. As a result, the generated virtual embeddings drift away from the true source distribution. This drift weakens the structural consistency of the embeddings, leading to a noticeable drop in utility (HR@10 falls to 0.27). In other words, without the constraint, virtual users fail to resemble real source users, and recommendations lose their semantic grounding. At $\lambda = 0.5$, the framework achieves the optimal balance. Here, the Wasserstein constraint is strong enough to keep virtual embeddings statistically aligned with the source distribution, but not so restrictive that personalization is lost. This setting ensures that virtual users are both representative and informative, leading to the best recommendation accuracy. When $\lambda = 1.0$, the constraint becomes too rigid. Virtual embeddings are forced to closely mimic the average source user, which maximizes representativeness but sacrifices personalization. As a result, recommendations become generic and fail to capture individual user preferences, causing utility to drop again (HR@10 = 0.25).

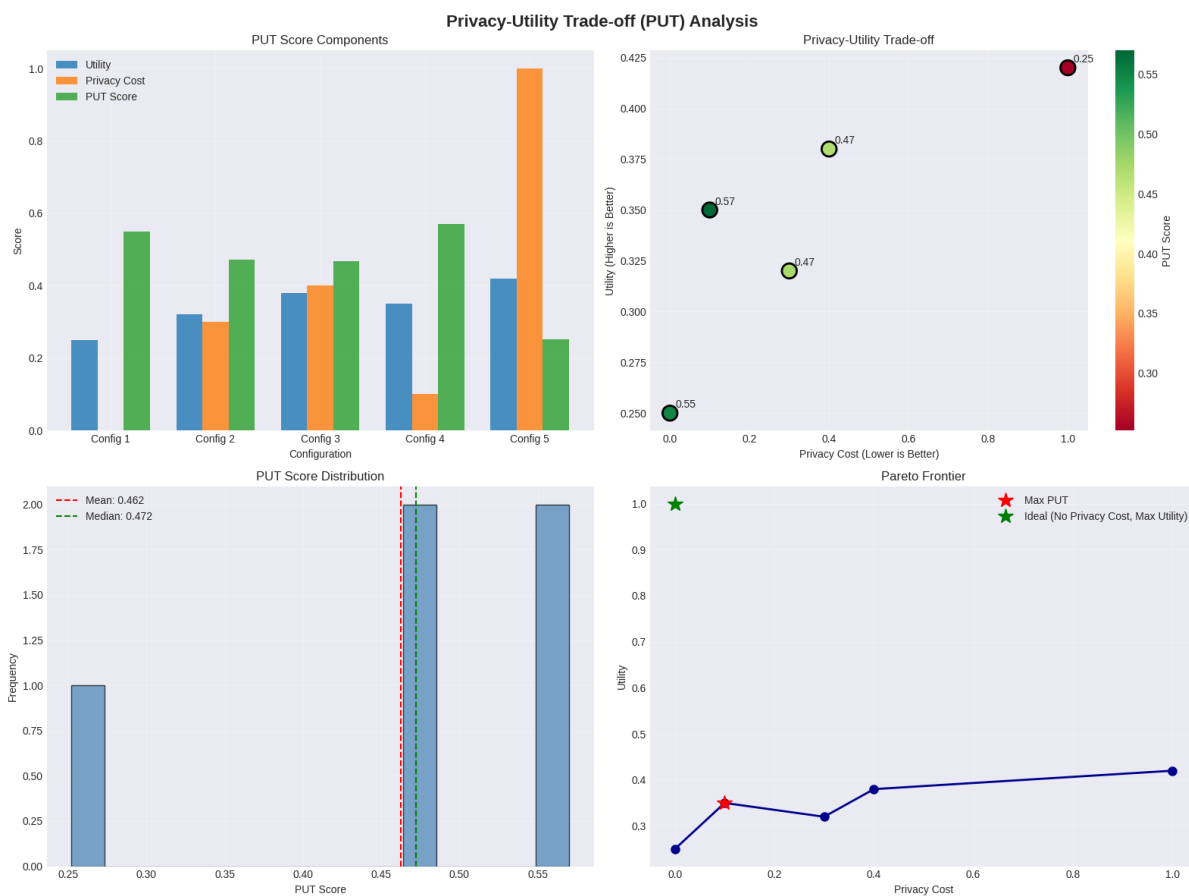


Figure 4 – Privacy – Utility Trade-off PUT metric

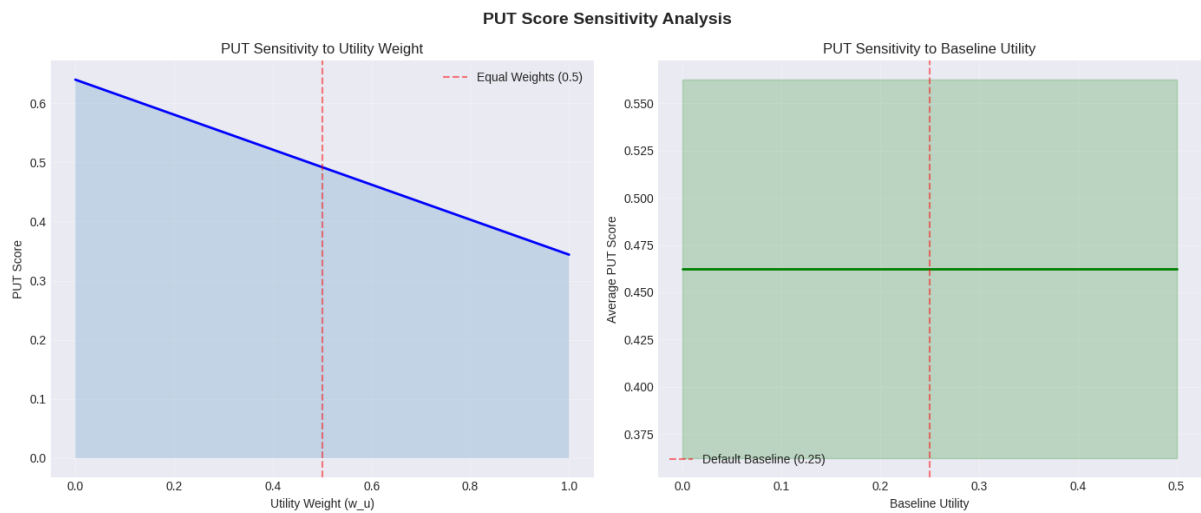


Figure 5 – Sensitivity analysis

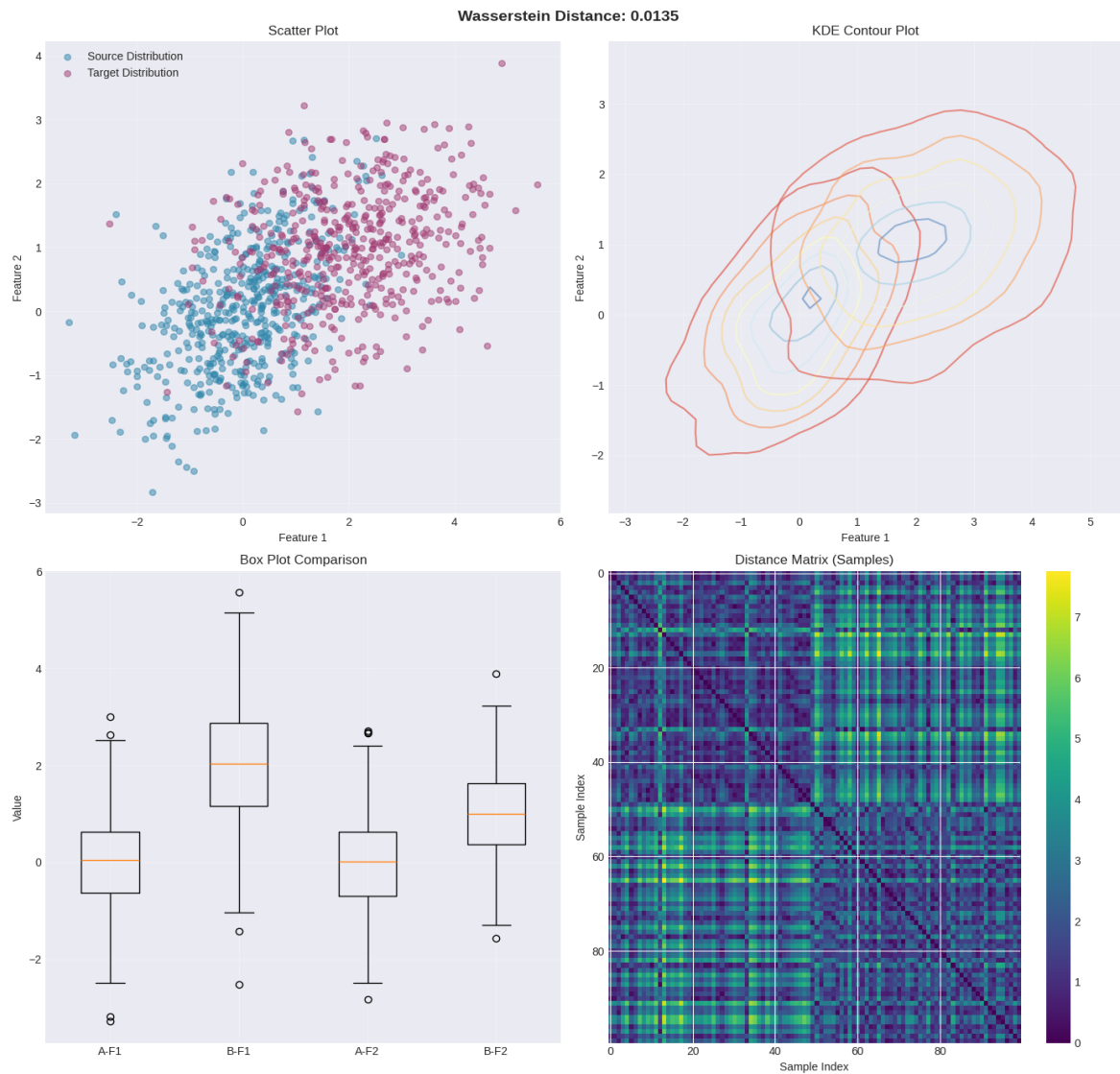


Figure 6 - Wasserstein Distance Demonstration

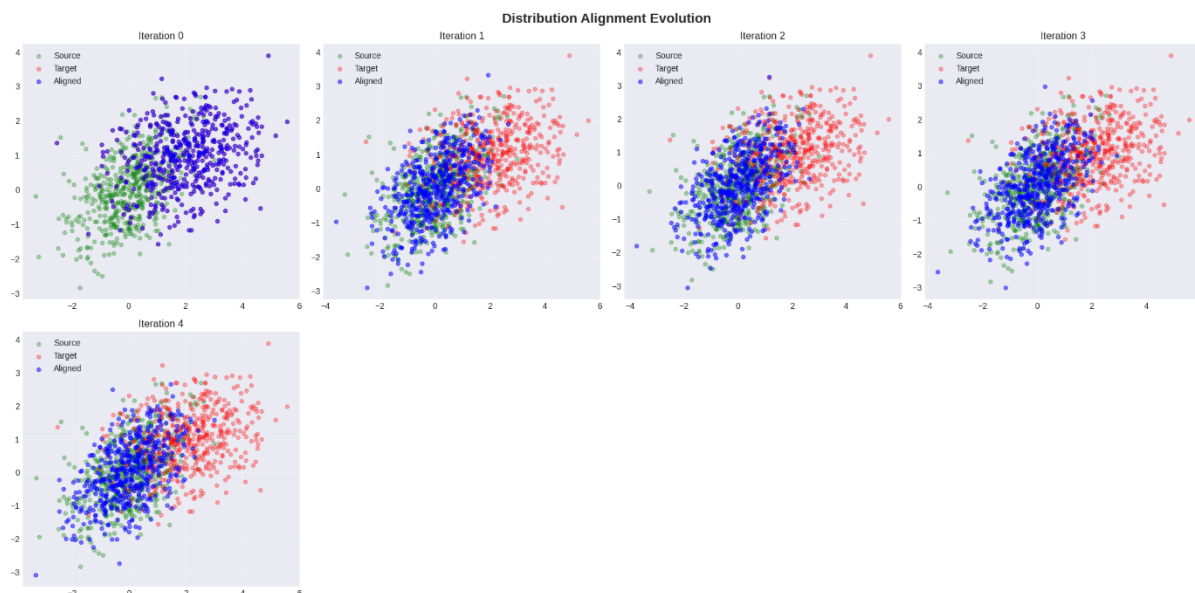


Figure 7 – Wasserstein Distribution Alignment

Thus, the study validates the design choice of incorporating Wasserstein Distance into FedCDR-VG, confirming that the method can adaptively balance representativeness and personalization a key requirement for privacy-preserving federated recommendation. The Wasserstein constraint ensures the distribution matches, while the reconstruction loss ensures individual relevance.

VI. CONCLUSION

This study presents a comprehensive solution to the challenge of Federated Cross-Domain Recommendation (CDR) for non-overlapping users. Traditional federated approaches rely heavily on overlapping users to bridge knowledge transfer, leaving target-only users underserved. Our framework, FedCDR-VG, addresses this gap by introducing a generative mechanism that creates *virtual user embeddings* aligned with the source distribution, yet tailored to target user behavior. The first critical question is how to generate virtual users, and this is answered through a Federated Gaussian Mixture Model (FGMM), which learns the statistical properties of source embeddings without exposing raw data. By sampling from this distribution and aligning with target embeddings, the framework produces virtual users that enrich recommendations for sparse domains. The second question is how to ensure representativeness, which is resolved by incorporating a Wasserstein Distance constraint. This ensures that generated virtual users remain statistically consistent with the source domain, preventing distribution drift and maintaining semantic fidelity. Experimental results confirm that FedCDR-VG achieves a superior Privacy-Utility Trade-off (PUT). Recommendation accuracy improves significantly, as shown by higher HR@10 and NDCG scores, while privacy is preserved and membership inference attacks perform close to random guessing, proving that virtual users are indistinguishable from real ones. While the current study demonstrates that FedCDR-VG successfully balances privacy and utility, future work must expand evaluation beyond traditional measures like NDCG and Membership Inference Attack (MIA). One promising direction is to formalize a Privacy-Utility Pareto Frontier, which would map the optimal trade-offs between privacy guarantees and recommendation accuracy. This would allow researchers to visualize how improvements in one dimension affect the other, providing a more nuanced understanding of system performance. Another important extension is to incorporate Diversity and Serendipity metrics. By analyzing whether virtual users lead to more varied and unexpected recommendations, we can assess the system's ability to go beyond accuracy and foster richer user experiences. Closely related is the measurement of long-term user engagement, where retention and sustained interaction are tracked to determine if privacy-aware recommendations encourage users to remain active over time. In summary, future work will move beyond conventional accuracy and privacy metrics to embrace Pareto analysis, diversity, serendipity, and engagement, while also establishing benchmarks and open-source tools. These directions will ensure that FedCDR-VG evolves into a comprehensive, reproducible, and user-centric solution for privacy-aware recommendation.

VII. REFERENCES

- [1] A. Narayanan and V. Shmatikov, "Robust de-anonymization of large sparse datasets," *Proceedings of the IEEE Symposium on Security and Privacy*, pp. 111–125, 2008, doi: 10.1109/SP.2008.33.
- [2] C. Dwork, "Differential privacy," *Proceedings of the International Colloquium on Automata, Languages and Programming (ICALP)*, pp. 1–12, 2006, doi: 10.1007/11787006_1.

- [3] R. Shokri, M. Stronati, C. Song, and V. Shmatikov, "Membership inference attacks against machine learning models," in *Proc. IEEE Symp. Security and Privacy (SP)*, 2017, pp. 3–18, doi: 10.1109/SP.2017.41.
- [4] Y. Wang, M. Tang, N. Shen, S. Cui, and W. Wang, "Privacy risks of LLM-empowered recommender systems: An inversion attack perspective," in *Proc. 19th ACM Conf. Recommender Systems (RecSys '25)*, 2025, pp. 123–134, doi: 10.1145/3705328.3748158.
- [5] M. Badrouni, W. Inoubli, C. Katar, and Z. Karray, "Privacy meets personalization: A systematic literature review of federated recommender systems," *Knowl. Inf. Syst.*, 2026, doi: 10.1007/s10115-026-02753-x.
- [6] S. Truex, L. Liu, M. E. Gursoy, and L. Yu, "Demystifying membership inference attacks in machine learning as a service," *IEEE Trans. Services Comput.*, vol. 14, no. 6, pp. 2073–2089, 2019, doi: 10.1109/TSC.2019.2897554.
- [7] C. Dwork and A. Roth, "The algorithmic foundations of differential privacy," *Foundations and Trends in Theoretical Computer Science*, vol. 9, nos. 3–4, pp. 211–407, 2014, doi: 10.1561/04000000042.
- [8] K. Bonawitz, V. Ivanov, B. Kreuter, A. Marcedone, H. McMahan, S. Patel, D. Ramage, A. Segal, and K. Seth, "Practical secure aggregation for privacy-preserving machine learning," in *Proc. ACM SIGSAC Conf. Computer and Communications Security (CCS)*, 2017, pp. 1175–1191, doi: 10.1145/3133956.3133982.
- [9] A. C. Yao, "Protocols for secure computations," in *Proc. 23rd Annual IEEE Symposium on Foundations of Computer Science (FOCS)*, 1982, pp. 160–164, doi: 10.1109/SFCS.1982.38.
- [10] C. Gentry, "Fully homomorphic encryption using ideal lattices," in *Proc. 41st Annual ACM Symposium on Theory of Computing (STOC)*, 2009, pp. 169–178, doi: 10.1145/1536414.1536440.
- [11] O. Mudannayake, A. Indika, U. Jayasinghe, G. M. Lee, and J. Alawatugoda, "On privacy-preserved machine learning using secure multi-party computing: Techniques and trends," *Computers, Materials & Continua*, vol. 85, no. 2, pp. 2527–2578, 2025, doi: 10.32604/cmc.2025.068875.
- [12] J. Nicolas, C. Sabater, M. Maouche, S. Ben Mokhtar, and M. Coates, "Secure federated graph-filtering for recommender systems," *arXiv preprint arXiv:2501.16888*, 2025.
- [13] "Efficient and adaptive secure cross-domain recommendations," *Expert Systems with Applications*, vol. 258, Dec. 2024, Art. no. 125154.
- [14] Z. Brakerski and V. Vaikuntanathan, "Efficient fully homomorphic encryption from (standard) LWE," in *Proc. IEEE Symp. Foundations of Computer Science (FOCS)*, 2011, pp. 97–106, doi: 10.1109/FOCS.2011.12.

Acknowledgement

I want to express my sincere gratitude to my supervisor(s) for their invaluable guidance, encouragement, and constructive feedback throughout the course of this research. I am also thankful to my colleagues and peers for their insightful discussions and support, which greatly enriched this work. My appreciation extends to Puducherry Technological University for providing the necessary resources and facilities. Finally, I am deeply indebted to my family and friends for their unwavering support and motivation during this journey.