

基于差分隐私的通信高效联邦推荐方法*

薛大暄^{1,2,3}, 杜宜霏^{1,2,3}, 陈红^{1,2,3}, 李翠平^{1,2,3}

¹(数据工程与知识工程教育部重点实验室(中国人民大学), 北京 100872)

²(数据库与商务智能教育部工程研究中心(中国人民大学), 北京 100872)

³(中国人民大学 信息学院, 北京 100872)

通信作者: 陈红, E-mail: chong@ruc.edu.cn



摘要: 推荐系统已成为大数据时代缓解信息过载问题的关键技术, 广泛应用于电子商务等领域, 但传统的集中式数据收集方式存在用户隐私泄露的风险. 联邦学习允许多个数据持有者在不共享用户原始数据的情况下进行联合训练以保护数据隐私, 联邦推荐系统也受到工业界和学术界的广泛关注. 现有的联邦推荐算法将推荐系统的建模过程置于分布式环境中, 有效避免了用户敏感信息在中心服务器上的集中存储, 但仍存在隐私泄露和通信成本高的问题. 针对该问题, 提出一种基于差分隐私的通信高效联邦推荐算法. 该算法设计一种通用的子模型选择策略, 通过在客户端采用随机响应机制加强对用户交互数据的隐私保护, 并在服务器端采用最大似然估计的方法估计物品的真实交互频率来优化子模型的选择, 实现用户隐私保护与模型效用之间的有效平衡. 该算法不仅适用于矩阵分解推荐模型, 还可扩展应用于深度学习推荐模型, 在不同建模场景下均表现出较高的灵活性和适用性. 此外, 为进一步降低通信开销, 针对深度学习模型复杂结构和庞大参数导致的通信负担, 提出全局模型结构化划分策略, 并为浅层网络和深层网络制定差异化的优化策略, 有效降低了通信开销. 理论分析表明该方法满足差分隐私性质. 在真实数据集上的实验结果表明, 该方法在不显著降低模型可用性的前提下, 保障了用户数据的隐私安全, 同时大幅提高了联邦推荐中的通信效率.

关键词: 联邦推荐; 差分隐私; 随机响应; 深度学习

中图法分类号: TP311

中文引用格式: 薛大暄, 杜宜霏, 陈红, 李翠平. 基于差分隐私的通信高效联邦推荐方法. 软件学报, 2026, 37(6): 2584–2606. <http://www.jos.org.cn/1000-9825/7550.htm>

英文引用格式: Xue DX, Du YF, Chen H, Li CP. Communication-efficient Federated Recommendation Method with Differential Privacy. Ruan Jian Xue Bao/Journal of Software, 2026, 37(6): 2584–2606 (in Chinese). <http://www.jos.org.cn/1000-9825/7550.htm>

Communication-efficient Federated Recommendation Method with Differential Privacy

XUE Da-Xuan^{1,2,3}, DU Yi-Fei^{1,2,3}, CHEN Hong^{1,2,3}, LI Cui-Ping^{1,2,3}

¹(Key Laboratory of Data Engineering and Knowledge Engineering of the Ministry of Education (Renmin University of China), Beijing 100872, China)

²(Engineering Research Center for Database and Business Intelligence of the Ministry of Education (Renmin University of China), Beijing 100872, China)

³(School of Information, Renmin University of China, Beijing 100872, China)

Abstract: Recommendation systems have become a key technology in mitigating information overload in the era of big data, with widespread applications in E-commerce and other fields. However, traditional centralized data collection methods expose significant risks of user privacy leakage. Federated learning enables collaborative model training across multiple data holders without the need to share raw user data, thus protecting privacy. Federated recommendation systems have gained considerable attention from both academia and industry.

* 基金项目: 国家重点研发计划 (2023YFB4503600); 国家自然科学基金 (U23A20299, U24B20144, 62172424, 62276270, 62322214)
收稿时间: 2025-01-09; 修改时间: 2025-04-30; 采用时间: 2025-09-04; jos 在线出版时间: 2026-02-11
CNKI 网络首发时间: 2026-02-12

Existing federated recommendation algorithms place the model training process in a distributed environment, effectively avoiding the centralized storage of sensitive user data on a single server. However, these approaches still face challenges related to privacy leakage and high communication costs. To address these issues, this study proposes a communication-efficient federated recommendation algorithm based on differential privacy. The algorithm introduces a general sub-model selection strategy that strengthens privacy protection of user interaction data on the client side through a randomized response mechanism. On the server side, it employs maximum likelihood estimation to infer the true interaction frequencies of items and optimize the sub-model selection process. This strategy achieves an effective balance between privacy protection and model utility. The proposed algorithm is applicable not only to matrix factorization-based recommendation models but also to deep learning-based models, demonstrating high flexibility and adaptability across various recommendation scenarios. Furthermore, to reduce communication overhead, a global model partitioning strategy is proposed to address the complex structures and large parameter sizes of deep learning models. Differentiated optimization strategies are applied to shallow and deep networks to effectively mitigate communication costs. Theoretical analysis shows that the method satisfies differential privacy, while experimental results on real-world datasets demonstrate that the proposed approach preserves user data privacy without significantly compromising model utility, while substantially improving communication efficiency in federated recommendation systems.

Key words: federated recommendation; differential privacy (DP); randomized response; deep learning

1 引言

随着数字技术的迅速发展,推荐系统作为应对信息过载问题、提升用户体验和决策效率的关键技术得到了广泛应用^[1].通过对用户历史行为、偏好特征及其上下文信息的分析,推荐系统能够智能地为用户提供个性化的信息或产品推荐,广泛应用于电子商务^[2]、社交网络和内容服务等领域.然而,推荐系统的高效运作高度依赖于对用户个人数据的深入挖掘,包括用户的基本属性、社交网络关系和行为数据等敏感信息.这种对数据的依赖在提高推荐准确性的同时,也带来了显著的隐私泄露风险.例如,某些攻击者可能通过分析用户的推荐记录,推测出该用户的敏感行为或偏好,甚至暴露出用户的身份信息,从而导致隐私数据被滥用的严重后果.这一隐私风险揭示了传统中心化数据收集与处理模式的安全隐患^[3,4].联邦学习作为一种保护用户隐私的分布式框架被引入推荐系统领域^[5-9],允许多个数据持有者在不直接共享用户数据的前提下,通过在用户设备上本地训练模型,并交换模型中间参数进行联合训练,在一定程度上兼顾了用户隐私保护与推荐服务的有效性,其通用架构如图1所示.然而,尽管联邦学习通过参数共享在一定程度上缓解了隐私泄露风险,其仍无法完全规避隐私攻击的威胁.此外,联邦学习的高效性依赖于多轮通信交互,而在大规模数据场景或资源受限设备中,这种频繁交互已面临通信成本高和模型复杂带来的严峻挑战^[10,11].更重要的是,为增强隐私保护而引入的差分隐私(differential privacy, DP)技术虽然通过噪声注入提升了隐私保障,但同时显著增加了通信开销,进一步加剧隐私保护与模型效用之间的不平衡问题.

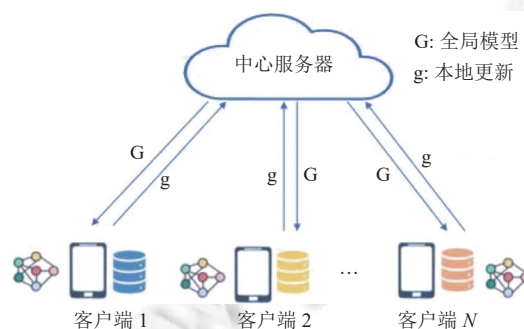


图1 联邦学习通用架构

一方面,虽然联邦学习框架在理论上为原始数据的隐私保护提供了有力保障,避免了直接传输用户数据.但研究表明,攻击者仍可能通过分析共享的模型参数或梯度信息来推测用户的敏感数据^[12,13].例如,文献[14]指出,通过对联邦学习过程中传输的梯度信息进行推测分析,攻击者可以重建部分用户的兴趣爱好或身份信息.为增强联邦学习中的隐私保护,差分隐私技术被引入该框架,通过在模型输出中加入噪声提供了数学上的隐私保证,从而防

止攻击者通过分析模型更新来推测个体数据. 文献 [15] 采用本地差分隐私技术实现推荐系统中用户级别的保护, 即对用户的评分和评分项目均进行扰动, 避免原始数据泄露. 除差分隐私外, 联邦推荐系统中还广泛采用安全多方计算、同态加密和秘密共享等密码学方法来加强隐私保护. 例如, 安全多方计算支持多方用户在不泄露私有数据的前提下完成联合计算, 同态加密允许在加密数据上执行推荐模型的前向与反向传播操作, 而秘密共享则通过参数分片传输提升抗泄露能力. 尽管上述方法在理论上提供了更强的隐私保障, 但其引入的建模复杂性和信息扰动在实践中可能影响模型训练过程的稳定性和精度表现. 特别是在数据量有限或交互行为稀疏的推荐场景中, 隐私机制所带来的性能退化问题更加突出. 这对联邦推荐系统的推荐准确度提出了更高的优化需求, 也进一步突显出在保障隐私的同时提升模型效能的重要性.

另一方面, 联邦推荐系统还面临着通信效率上的重大挑战. 由于联邦学习要求客户端与服务器频繁交互, 每轮训练需要传输大量的全局模型参数进行学习, 这在面对大规模学习任务和资源受限的移动设备时显得非常低效^[16]. 特别是在推荐系统的应用场景中, 每个用户的数据通常包含多维度、复杂的行为信息, 而推荐模型往往需要处理庞大的物品集. 为了在用户设备上进行有效的个性化推荐, 联邦推荐系统通常采用从简单到复杂的模型体系, 例如矩阵分解模型^[17]和深度学习模型^[18]. 矩阵分解模型因其参数量相对较小、通信成本较低, 成为联邦推荐系统中优化通信效率的基础方法^[19]. 然而, 随着用户数量和交互数据规模的增加, 即使是矩阵分解模型也可能面临通信成本快速上升的问题, 尤其是在多用户协作频繁或矩阵规模较大的场景中. 相比之下, 深度学习模型凭借其强大的特征提取和表达能力, 能够处理更复杂的数据模式, 在提升推荐精度方面具有明显优势. 然而, 其参数规模远超矩阵分解模型, 训练过程中的通信成本和隐私保护问题更为突出. 这些模型含有大量参数, 且在每轮训练迭代中, 客户端需要将更新后的模型参数上传至服务器. 而在服务器端, 必须在汇总多个客户端上传的参数后再将更新的全局模型下发至各客户端. 以 MovieLens 1M 为例, 若包含两层隐藏层的多层感知机作为推荐模型, 其参数量可达约 3–5 百万级别, 假设采用 32-bit 浮点编码, 每轮训练需在客户端与服务器之间传输约 12–20 MB 的参数数据. 若每轮迭代时有 100 个客户端, 整体传输数据可达 GB 级别. 对于这种频繁的模型更新传输进一步放大了通信负担, 最终导致系统整体的运行效率下降, 并可能使系统在实际应用中不可扩展. 这种情况对于推广联邦学习技术在推荐系统中的应用构成了实质性的障碍.

为解决上述问题, 本文提出一种新的解决策略, 即基于差分隐私的通信高效联邦推荐算法. 该算法的核心思想是设计一种通用“子模型”选择机制. 在推荐模型的训练和通信过程中优化数据传输量, 降低通信成本, 通过结合差分隐私技术, 实现在矩阵分解和深度学习模型中对通信效率与用户隐私保护的双重优化. 其主要贡献如下.

(1) 针对联邦推荐系统中通信开销高和用户交互数据隐私易泄露的问题, 提出基于差分隐私的子模型选择算法 DP-SUB. 该算法允许客户端仅下载和更新所需的模型部分, 并通过随机响应机制加强了用户交互数据的隐私保护, 在降低通信成本的同时, 有效保障了用户数据的隐私安全.

(2) 针对服务器难以准确推断客户端需求的问题, 提出基于最大似然估计的优化算法. 通过估计物品的真实交互频率来优化子模型的选择, 确保在保护用户交互数据隐私的前提下, 服务器能够高效且准确地推断并下发符合客户端需求的子模型, 实现了用户隐私保护与模型效用之间的有效平衡.

(3) 针对深度学习的联邦推荐模型中, 因结构复杂导致通信负担加重的问题, 提出基于差分隐私的通信高效联邦深度学习推荐算法 Priv-FedNCF-Sub-Compress. 通过对全局模型进行结构化划分并对浅层网络和深层网络定制不同的优化策略, 有效降低了通信开销. 同时, 采用满足差分隐私保护的梯度扰动策略和基于 FedSubAvg 算法的子模型聚合优化技术, 进一步平衡了隐私保护、通信效率和推荐性能之间的关系.

本文第 1 节介绍联邦推荐的研究现状. 第 2 节介绍联邦推荐中的相关工作. 第 3 节介绍本文所需的基础知识, 包括差分隐私和联邦推荐. 第 4 节介绍基于差分隐私的子模型选择策略, 包括基于随机响应的交互数据扰动策略和基于最大似然函数的真实频率估计. 第 5 节分别介绍本文构建的基于子模型的联邦矩阵分解推荐算法和基于差分隐私的通信高效联邦深度学习推荐算法. 需要说明的是, 两种推荐方法在应用场景上具有互补性. 矩阵分解推荐算法以其较低的通信成本和轻量级模型设计, 适用于对隐私保护敏感但资源有限的场景. 而深度学习推荐算法则凭借其在提取复杂特征和处理大规模数据方面的优势, 适用于对推荐精度和数据模式要求较高的场景. 第 6 节通

过对比实验验证了所提方法的有效性. 最后总结全文.

2 相关工作

2.1 联邦推荐中的模型方法

在传统的集中式推荐系统中, 所有用户的个人数据都被集中存储在中心化服务器上, 推荐算法能够直接访问这些数据, 进而提供个性化的推荐服务. 尽管这种方式提高了推荐效率, 但集中处理大量敏感数据极易导致用户隐私泄露^[20]. 为了解决这一问题, 研究者提出了基于联邦学习框架的推荐算法. 该框架通过将用户数据保存在本地设备上, 采用协作学习的方式, 在不直接访问用户原始数据的情况下完成模型训练, 有效降低了隐私泄露的风险.

现有的联邦推荐算法主要分为两类: 基于矩阵分解^[21]和基于深度学习^[22]的推荐算法. 基于矩阵分解的联邦推荐算法是目前应用最广泛的方法, 其核心思想是将集中式数据收集方式转变为分布式数据处理方式, 以在不直接接触用户数据的情况下实现个性化推荐. 各参与方通过协同学习, 利用各自的用户行为数据, 共同构建一个全局推荐模型. 文献[23]将联邦学习引入推荐系统, 提出了联邦协同过滤 (federated collaborative filtering, FCF) 推荐算法. 该算法通过本地保存用户交互数据, 更新用户特征向量, 服务器端则通过聚合客户端上传的梯度来更新物品特征向量. 需要注意的是, FCF 在处理隐式反馈数据时, 将未评分的物品视为未交互项, 这种假设可能引入偏差, 影响模型的准确性. 此外, FCF 并未针对通信效率进行优化, 每轮仍需传输完整的物品特征向量矩阵, 导致通信开销较高. 文献[24]提出了 FedRec 算法, 专注于显式反馈数据. 该算法通过随机采样未交互数据并分配虚拟评分的方法, 在提高模型效用的同时有效降低了计算和通信成本. 然而, 该算法对中间参数泄露的隐私风险估计不足, 且降低通信成本的效果尚未得到实证验证. 尽管基于深度学习的联邦推荐算法面临着客户端存储和计算能力有限的问题, 难以支持大规模神经网络模型的训练, 但由于其在推荐准确性和个性化方面的潜力, 仍然引起了广泛的关注. 文献[25]提出了联邦视频推荐框架 JointRec, 使用卷积神经网络 (convolutional neural network, CNN) 提取用户属性、视频属性和评论特征, 以完成视频推荐. 由于模型结构复杂, 导致通信开销过高, 并且未充分考虑可能出现的梯度泄露风险. 文献[26]提出了 FedDSR 方法, 结合深度强化学习和联邦学习, 利用课程学习指导训练过程, 并通过相似性聚合算法提升上传本地参数的质量. 但在 FedDSR 中, 各方需要频繁地上传和下载本地模型参数, 尤其是在大规模数据集的情况下, 可能导致通信频率和数据量较大, 从而增加了通信开销. 文献[27]提出了基于联邦分布式深度确定性策略梯度方法 FD3PG, 通过部分可观察的马尔可夫决策过程 POMDP 来优化多层边缘-云网络中的内容传输延迟、缓存替换和带宽分配策略. 然而频繁的参数更新通信、较高的带宽消耗、同步问题带来的延迟以及隐私保护带来了额外通信开销. 文献[28]提出的 Uni-FedRec 框架通过聚类方法和注意力机制生成共享的基础嵌入, 有效防止了用户隐私泄露, 但全局模型传输仍然存在通信效率低下的问题.

2.2 联邦推荐中的隐私保护机制

为了进一步提高联邦推荐中的隐私保护, 文献[29]提出了 FedRec++, 通过在客户端实施隐私感知去噪, 显著改善了推荐质量. 为了解决隐私与公平性问题, 文献[30]提出了公平感知的联邦矩阵分解框架 F2MF, 结合差分隐私技术与联邦学习系统, 确保在不暴露用户敏感群组特征的情况下实现推荐系统的公平性. 值得注意的是, 差分隐私引入的噪声对推荐精度产生了影响, 尤其是在数据较稀疏的场景中. 文献[31]提出了 MetaMF 框架, 通过服务器上的元网络生成私有项目嵌入, 有效提升了推荐系统的性能和效率, 并降低了用户设备的计算负担. 但是 MetaMF 在处理冷启动用户时效果不佳, 且依然存在潜在的隐私泄露风险. 异构协同过滤算法 FCMF^[32]结合了用户反馈和同态加密技术, 确保在异构场景下的隐私保护. 文献[33]提出了一种基于秘密共享的联邦矩阵分解方法, 通过将模型参数随机分割并利用秘密共享技术在用户和服务器间传输, 从而保护用户隐私. 文献[8]提出的 FedFast 算法将 GMF^[34]应用于联邦学习, 通过客户端采样提高收敛速度, 同时采用安全聚合技术保护上传梯度的隐私. 尽管 FedFast 在推荐效果上有所提升, 但复杂的加密和解密过程增加了计算和通信的负担, 特别是在资源受限的设备上. 类似地, FedGNN 框架^[35]引入第三方服务器并采用同态加密技术匿名传递邻居用户特征向量^[36], 进一步加强了隐私保护, 但同态加密的应用也增加了系统的计算和通信负担. 文献[37]提出的 FedPerGNN 框架通

过隐私保护的模型更新方法和图扩展协议,充分利用去中心化的图数据进行个性化推荐,但高阶图信息的利用效率和通信开销问题仍有待进一步解决.总体而言,现有联邦推荐系统中的隐私保护方法可大致归为 3 类:其一,基于差分隐私的扰动机制,通过在用户数据或模型更新中引入噪声实现个体级或群体级隐私保护;其二,基于加密机制的隐私保护方法,包括同态加密与秘密共享技术有效防止中间值泄露;其三,混合机制,结合结构设计与保护策略,如 MetaMF 在服务器端生成私有项目嵌入, FedPerGNN 则融合图扩展与隐私更新协议实现去中心化推荐.相比加密机制带来的计算与通信负担,差分隐私具备实现成本低、可调节性强等优势,尤其适合资源受限设备场景.

2.3 联邦推荐中的通信优化技术

为应对高频率模型同步所带来的带宽消耗与系统负担,文献[38]提出了深度梯度压缩机制,通过梯度稀疏化、动量修正、局部裁剪等技术实现高达 600 倍的梯度压缩比,在不降低模型精度的前提下显著减少通信负担.文献[39]进一步提出 FetchSGD 算法,采用可合并的 Count Sketch 结构压缩模型更新,并将动量与误差累积迁移至服务端执行,使得在客户端参与率低或模型大规模时仍可实现高效、鲁棒的通信压缩.在模型结构优化方面,文献[40]提出 SlimFL 方法,将可调宽度的可瘦身神经网络与联邦学习结合,并引入叠加编码与叠加训练策略,以实现不同模型宽度配置下的高效参数共享.文献[41]则提出 FedSVD 方法,通过局部梯度压缩与全局特征共享策略,在保证精度无损的前提下完成亿级规模数据上的联邦奇异值分解任务,为超大规模联邦推荐提供了通信可控的建模方案.客户端选择与调度策略也是降低通信成本的重要路径.文献[42]提出 HiCS-FL 框架,基于客户端上传的输出层更新估计其数据异质性,并通过层次化聚类方式筛选代表性客户端参与模型更新,减少冗余通信.文献[43]针对联邦学习中客户端响应时延不一致的问题,引入一种缓冲式异步聚合机制,通过仅对先到达的一部分客户端更新进行分批聚合,避免等待全部设备完成通信,从而在保证训练稳定性的同时有效提升系统效率.在编码与信息传输机制方面,文献[44]提出 GossipFL 框架,通过稀疏化通信与去中心化结构设计,使每个客户端仅与一个邻居通信并交换高度稀疏化的模型参数,同时构造自适应带权通信矩阵,缓解中心化瓶颈.文献[45]的 FedBoost 方法则从模型训练机制入手,采用集成学习策略,在客户端训练多个轻量级预测子模型,服务器端再进行加权集成,有效规避了训练完整大模型所需的高通信成本.此外,文献[46]提出 FAST 框架,通过自适应数据采样与本地训练联合优化机制,引导客户端根据本地类别重构训练数据集,在给定时间预算下加速模型收敛,并在保证通信资源受限下取得良好的性能表现.总体而言,现有通信优化方法大多聚焦于深度神经网络结构的通信压缩与局部训练策略,其核心在于通过模型剪枝、梯度稀疏化或客户端调度等方式缓解大规模模型训练过程中的带宽瓶颈.然而,这些方法普遍建立在复杂模型结构基础上,难以适应如矩阵分解等浅层轻量模型的特点,且部分方法在非 IID 数据分布下稳定性不高.

综上所述,现有的联邦推荐系统通过矩阵分解和深度学习等方法在隐私保护与推荐性能方面取得了一定进展,但仍面临多重挑战.一方面,矩阵分解方法虽然计算复杂度较低,但在隐式反馈数据处理和通信效率优化方面存在不足;另一方面,基于深度学习的方法在推荐准确性和个性化方面表现出色,但由于模型复杂性和参数量大,带来了显著的通信和计算负担.针对这些问题,进一步研究如何在保护隐私的同时优化通信效率、提升模型性能是联邦推荐系统未来的重要研究方向.

3 基础知识

3.1 差分隐私

差分隐私^[47]是一种数学机制,旨在在数据分析中保障个体隐私,通过向输出结果中引入随机噪声,减少单个个体对分析结果的影响,防止攻击者推断出个体信息.该机制广泛应用于大数据分析和机器学习中,尤其是在处理敏感数据时,能够有效提供隐私保护.具体定义如下.

定义 1. 差分隐私. 设 M 为随机算法,若对于任意两个相邻数据集 D_1 和 D_2 (即仅相差一个元素) 以及任意输出子集 S , 如果概率 P 满足公式 (1):

$$P[M(D_1) \in S] \leq e^\epsilon P[M(D_2) \in S] \quad (1)$$

则称 M 是 ϵ -差分隐私的. 其中 ϵ 为隐私预算, 表示隐私保护强度, 值越小则隐私保护越强.

差分隐私根据系统架构的不同, 可以分为中心差分隐私和本地差分隐私. 前者假设用户数据集中存储在一个可信的中央服务器上, 而后者要求用户在本地进行数据扰动, 再上传扰动后的数据到服务器, 避免将数据暴露给中心化机构.

定义 2. 中心差分隐私. 指用户数据集中存储在可信服务器中, 服务器通过向分析结果中添加噪声保证差分隐私. 常用的噪声机制包括高斯机制和拉普拉斯机制, 高斯机制通过向查询结果中添加正态分布的噪声来实现隐私保护, 在一些需要更宽松隐私预算的场景中具有更好的灵活性. 拉普拉斯机制通常适用于满足 ϵ -差分隐私的场景, 通过对实值查询结果添加服从拉普拉斯分布的噪声来实现隐私保护. 噪声幅度 b 与查询函数的敏感度 Δf 成正比:

$$Lap(b), b = \frac{\Delta f}{\epsilon} \quad (2)$$

定义 3. 本地差分隐私. 要求用户在本地对数据添加噪声后再上传^[48]. 若对于用户的任意两个不同输入 x_1 和 x_2 , 本地算法 M 满足:

$$P[M(x_1) = y] \leq e^\epsilon P[M(x_2) = y] \quad (3)$$

则称 M 具有本地 ϵ -差分隐私. 本地差分隐私通过本地扰动数据^[49], 确保用户上传数据前隐私已被保护.

定义 4. 敏感度. 敏感度表示对于任意两个相邻数据集 D_1 和 D_2 , 函数 f 的最大变化量 Δf 公式为:

$$\Delta f = \max_{D_1, D_2} \|f(D_1) - f(D_2)\| \quad (4)$$

敏感度衡量查询结果对单个数据点的影响, 是决定噪声大小的关键因素.

定理 1. 组合性. 对于数据集 D 和 n 个满足差分隐私算法的 M_i , 若每个算法 M_i 均满足 ϵ_i -差分隐私, 则这些算法组合后的整体隐私预算为 $\sum_{i=1}^n \epsilon_i$. 这意味着多次差分隐私查询在同一数据集上执行时, 隐私预算会逐次累加.

定理 2. 并行性. 将数据集 D 分成 k 个互不相交的子集 D_i , 并对每个子集分别应用满足 ϵ_i -差分隐私的算法 M_i , 则这些算法并行执行时, 整体隐私预算为 $\max(\epsilon_i)$. 这表明, 在不同数据子集上并行进行差分隐私查询时, 隐私预算不发生累积.

定理 3. 后处理性. 若算法 M 满足 ϵ -差分隐私, 则对其输出结果进行任意后处理操作, 隐私保护仍然有效. 即一旦差分隐私算法的输出生成, 无论对结果进行何种处理或转换, 差分隐私的保护强度不变.

3.2 联邦推荐

设有 N 个客户端 (用户设备), 每个客户端 i 拥有其本地的用户-物品交互数据集 D_i , 其中包含用户的点击、浏览、评分等行为记录. 为了保护用户隐私, 客户端的数据不会上传至服务器. 联邦推荐系统的目标是在隐私保护的前提下, 通过各客户端协同训练全局推荐模型 w , 并为所有客户端提供高质量的个性化推荐服务^[50]. 在联邦推荐系统中, 每个客户端在其本地数据 D_i 上进行推荐模型训练, 服务器依据各客户端数据集的大小对模型参数 w_i 进行加权, 生成全局推荐模型 w . 具体模型聚合过程可表示为:

$$w = \sum_{i=1}^N p_i w_i, p_i = \frac{|D_i|}{\sum_{i=1}^N |D_i|} \quad (5)$$

其中, p_i 表示客户端 i 在全局模型中的权重, 依据客户端本地数据集大小进行分配. 较大的数据集将在模型更新中占据更大的权重, 从而保证模型对不同客户端数据的适应性. 全局优化问题可以表示为:

$$w = \operatorname{argmin} \sum_{i=1}^N p_i f_i(w, D_i) \quad (6)$$

其中, $f_i(w, D_i)$ 是客户端 i 在其本地数据集 D_i 上的损失函数, 通常用于衡量模型对用户-物品交互行为的预测误差. 推

荐系统中的损失函数设计多样,可基于用户对物品的评分预测或点击行为进行优化.联邦推荐系统的工作流程如下.

- (1) 本地训练: 每个客户端 i 在本地数据集 D_i 上训练推荐模型并更新参数 w_i .
- (2) 上传参数: 本地模型训练完成后, 客户端将更新的模型参数 w_i 上传至中央服务器.
- (3) 模型聚合: 服务器端对各客户端上传的模型参数 w_i 按照数据集大小 D_i 加权, 生成全局模型 w .
- (4) 下发模型: 服务器将全局模型 w 下发至各客户端, 继续下一轮本地训练.
- (5) 迭代优化联邦推荐系统多次迭代, 持续优化全局模型, 直到模型收敛或推荐效果达到预期

4 基于差分隐私的子模型选择策略

在横向联邦学习中^[51], 服务器被假设为诚实但好奇的, 推荐系统需在隐私保护与通信效率之间实现平衡. 在传统架构中, 客户端需下载完整的全局模型并上传更新参数, 物品数量庞大时通信开销显著. 针对这一问题, 本文提出了一种基于差分隐私的通信高效联邦推荐方法, 该方法设计了一种通用的子模型选择策略, 使客户端仅下载与其交互相关的子矩阵, 以降低通信负担. 然而, 此方法面临的关键挑战在于子模型的位置与用户交互数据直接关联, 若直接向服务器透露所需子模型的位置, 可能导致用户偏好信息泄露. 用户的交互行为 (如购买、点击等) 往往包含敏感隐私数据, 可能反映其健康状况、个人兴趣等. 为此, 本文提出基于差分隐私的子模型选择算法 DP-SUB, 在保障用户隐私的前提下实现高效子模型选择, 有效降低通信成本.

图2展示了基于差分隐私的子模型选择算法 DP-SUB 框架. 在该框架中, 用户首先将物品的真实交互索引编码为二进制向量, 并在本地应用满足本地差分隐私 (local differential privacy, LDP) 约束的随机响应机制, 对交互数据进行扰动, 以确保用户真实交互行为的隐私性. 扰动后的交互向量由用户上传至中心服务器, 服务器将接收到的所有扰动向量进行聚合, 形成全局扰动交互矩阵. 基于该矩阵, 服务器计算物品的整体交互频率, 并利用最大似然估计方法对物品的真实交互频率进行复原. 基于频率筛选的策略并结合预设的频率阈值, 服务器筛选出高频交互向量作为下一轮训练所需的子模型, 并将其下发至客户端, 从而在隐私保护的约束下实现通信效率的提升. 在通信复杂度方面, 假设每轮训练中服务器需向客户端下发模型参数 $W \in \mathbb{R}^d$, 传统方法需传输完整模型, 通信复杂度为 $O(d)$. 在 DP-SUB 选择策略中, 仅传输高频交互向量的子模型参数 $W_{\text{sub}} \in \mathbb{R}^k$, 通信复杂度降低为 $O(k)$, 其中 $k \ll d$, 通信成本显著下降.

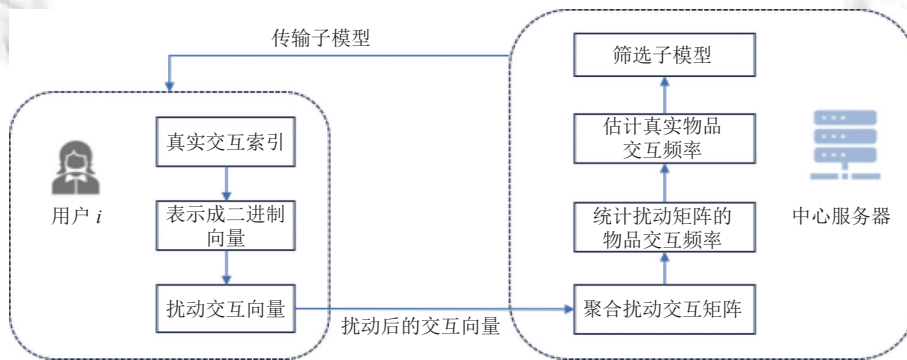


图2 基于差分隐私的子模型选择算法 DP-SUB 框架

4.1 基于随机响应的交互数据扰动策略

为确保用户在参与联邦学习推荐系统时的交互数据不被泄露, 本文在客户端侧引入本地差分隐私 (LDP) 机制. 这是一种在客户端对用户原始数据进行本地随机化处理的隐私保护机制, 其目标是在不依赖于中心服务器可信性的前提下, 防止攻击者通过观测上传数据推断用户的真实输入. 形式化地, 设随机化机制 $\mathcal{M}(\cdot)$ 作用于用户的原始输入 x , 若对于任意两个可能的输入 x, x' 以及任意输出结果 y , 均满足 $\Pr[\mathcal{M}(x) = y] \leq e^\epsilon \Pr[\mathcal{M}(x') = y]$, 则称机制 $\mathcal{M}(\cdot)$ 满足 ϵ -本地差分隐私 (ϵ -LDP), 其中 ϵ 为隐私预算参数, 用于刻画隐私保护与数据可用性之间的权衡. 具体

而言, 每个客户端首先将其原始交互矩阵 $S = [s_{ij}]$ 转化为二进制矩阵 $S' = [s'_{ij}]$, 其中 $s'_{ij} = 1$ 表示用户 i 与物品 j 存在交互, $s'_{ij} = 0$ 表示无交互. 该矩阵 S' 是用户的真实交互状态数据, 未经保护直接上传会造成隐私泄露. 为了对 S' 进行隐私保护, 本文应用差分隐私的随机扰动机制^[52,53]对其进行扰动. 具体地, 对于每个元素 s'_{ij} , 扰动后得到的结果 $\tilde{s}_{ij}$ 满足以下概率分布:

$$\tilde{s}_{ij} = \begin{cases} s'_{ij}, & \text{保持概率 } \frac{e^\epsilon}{e^\epsilon + 1} \\ 1 - s'_{ij}, & \text{翻转概率 } \frac{1}{e^\epsilon + 1} \end{cases} \quad (7)$$

其中, 参数 ϵ 为隐私预算, 其值越大, 隐私保护的强度越低, 但数据的真实性更高; 相反, 较小的 ϵ 值则提供更高的隐私保护, 然而会增加噪声的引入. 图 3 展示了这一扰动过程的示意图. 扰动后的交互数据矩阵 $\tilde{S} = [\tilde{s}_{ij}]$ 在保证用户隐私的前提下, 被上传至服务器. 该过程保证了每个用户的真实交互信息在传输过程中被有效“掩盖”, 即便攻击者获得了上传数据, 也无法准确推断用户的实际交互状态.

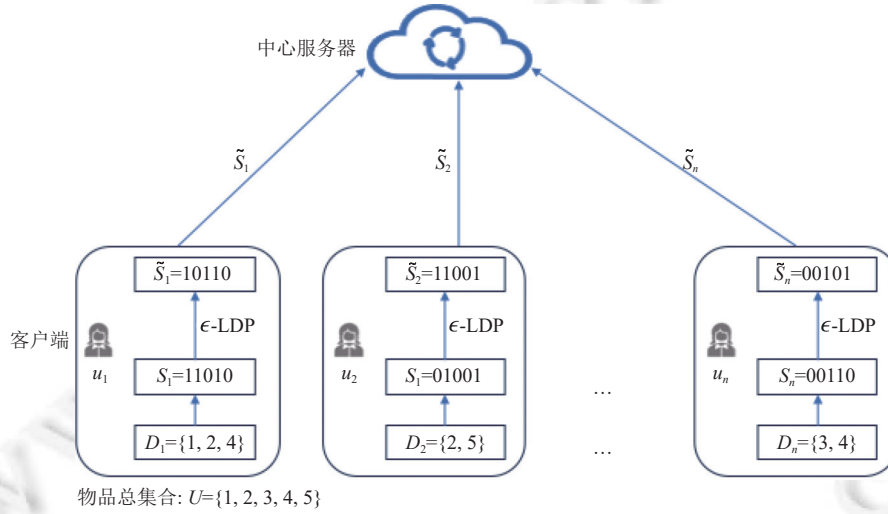


图 3 基于随机响应的交互数据扰动策略示意图

4.2 基于最大似然函数的真实频率估计

由于客户端上传的交互数据经过了差分隐私扰动, 直接利用这些数据进行推荐模型的训练将导致较大的偏差. 因此, 为了有效去除扰动带来的噪声干扰, 服务器端需要对扰动数据进行去噪处理, 以恢复物品的真实交互频率. 本文采用了最大似然估计 (maximum likelihood estimation, MLE) 方法^[54], 通过对扰动数据进行统计分析, 估计物品的真实交互概率 $\hat{p}_j$.

设服务器端观测到的扰动矩阵为 $\tilde{S} = [\tilde{s}_{ij}]$, 其中 $\tilde{s}_{ij}$ 表示用户 u_i 与物品 j 的扰动交互状态. 每个 $\tilde{s}_{ij}$ 值是由用户的真实交互状态 s'_{ij} 通过差分隐私扰动机制生成. 服务器只能观测到扰动后的交互频率 $\tilde{p}_j$, 即物品 j 的观测频率分布. 由随机响应扰动策略公式 (7) 可知, 当真实交互状态为 s'_{ij} 时, 无论 s'_{ij} 是 0 还是 1, 保持原值的概率为 $p_{00} = p_{11} = \frac{e^\epsilon}{e^\epsilon + 1}$, 而翻转的概率为 $p_{01} = p_{10} = \frac{1}{e^\epsilon + 1}$. 假设有 n 个用户参与本轮训练, 服务器收到扰动矩阵 $\tilde{S}$ 中第 j 列中数值为 1 的个数为 c , 则该列中数值为 0 的个数为 $n - c$. 根据最大似然估计原理, 构造对数似然函数 $\log(L) = c \times \log \Pr(s'_{ij} = 1) + (n - c) \times \log \Pr(s'_{ij} = 0)$, 对函数求导并令一阶导数为 0, 可得最大似然估计值. 具体公式为:

$$\hat{p}_j = \frac{p_{00} - 1}{p_{00} + p_{11} - 1} + \frac{\tilde{p}_j}{p_{00} + p_{11} - 1} \quad (8)$$

其中, $\tilde{p}_j$ 为物品 j 的观测频率, $p_{00} = p_{11}$ 为已知的保持概率, 其数值取决于隐私预算参数 ϵ . 由于对数似然函数二阶

导数为 0, 保证了解的唯一性和最大值存在性. 通过公式 (8), 服务器能够从扰动交互数据中准确恢复物品的真实交互频率 ρ_j , 在保证隐私保护的同时, 为推荐系统提供更加精确的偏好信息. 具体流程如图 4 所示, 服务器首先对所有客户端的扰动交互数据进行聚合, 形成整体扰动交互矩阵 $\tilde{S}$, 并基于该矩阵计算各物品 j 的观测交互频率 $\tilde{\rho}_j$. 随后, 利用最大似然估计公式 (8) 对每个物品的真实交互频率 ρ_j 进行恢复, 生成频率估计向量 $\hat{\rho} = [\hat{\rho}_1, \hat{\rho}_2, \dots, \hat{\rho}_m]$. 最后, 根据估算得出的真实交互频率筛选出交互频率超过预设阈值 p 的子模型向量, 为后续的推荐模型训练提供高质量的数据输入. 通过上述过程, 服务器能够在有效去除噪声影响的基础上, 精确恢复用户的偏好信息, 确保推荐系统在差分隐私保护下具备高精度的推荐效果.

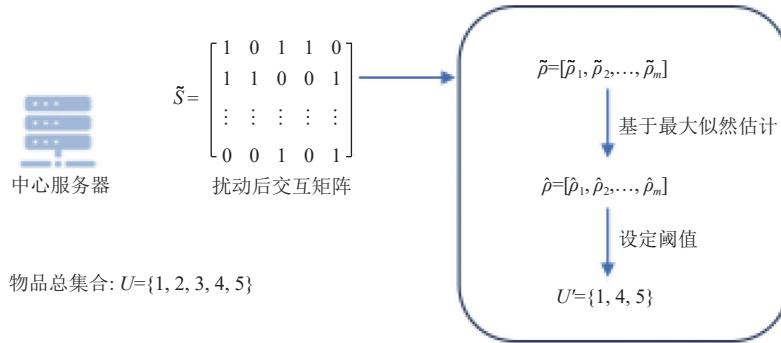


图 4 基于最大似然函数的真实频率估计示意图

4.3 DP-SUB 算法概述

该算法的目标是通过客户端和服务端端的协作, 在保护用户交互数据隐私的同时, 实现对物品交互频率的准确估计, 并优化子模型选择过程. 算法主要包含两部分: 客户端侧通过基于随机响应的策略对交互数据进行扰动, 并上传至服务器; 服务器侧则结合扰动后的数据, 采用最大似然估计方法对交互频率进行还原, 并完成子模型选择. 算法 1 为 DP-SUB 算法的具体步骤概述.

算法 1. DP-SUB 算法.

输入: 物品总集合 $U = \{1, 2, \dots, m\}$, 用户 u_i 的交互数据集 $D_i \subset U$, 隐私预算 ϵ_1 , 预设阈值 p ;
输出: 子模型集合 U' .

客户端:

1. 初始化客户端用户向量 S_i , 扰动向量 $\tilde{S}_i$
2. for $j = 1$ to m do
3. if U 中第 j 个值在 D_i 中: $s_{ij} = 1$, 否则 $s_{ij} = 0$
4. 根据 $\tilde{s}_{ij} = \begin{cases} s_{ij}, & \text{保持概率 } \frac{e^\epsilon}{e^\epsilon + 1} \\ 1 - s_{ij}, & \text{翻转概率 } \frac{1}{e^\epsilon + 1} \end{cases}$ 扰动交互向量
5. end for
6. 上传扰动向量 $\tilde{S}_i$ 至中心服务器

服务器端:

1. 初始化服务器端扰动频率 $\tilde{\rho}$, 真实频率 $\hat{\rho}$, 子模型集合 U'
2. 计算扰动概率 $p_{00} = p_{11} = \Pr(\tilde{s}_{ij} = s_{ij}) = \frac{e^\epsilon}{e^\epsilon + 1}$, 聚合形成扰动交互矩阵 $\tilde{S}$
3. for $j = 1$ to m do
4. 计算扰动物品交互频率 $\tilde{\rho}_j$

5. 根据 $\hat{p}_j = \frac{p_{00}-1}{p_{00}+p_{11}-1} + \frac{\tilde{p}_j}{p_{00}+p_{11}-1}$ 估计真实交互频率 $\hat{p}$
6. end for
7. 根据阈值 p 和真实交互频率 $\hat{p}$, 筛选出交互频率超过阈值的子向量, 并将其 ID 添加至子模型集合 U'
8. return U'

5 基于不同场景下的联邦推荐模型适配设计

矩阵分解方法因其结构简洁与通信代价低, 长期作为联邦推荐系统的基础建模手段, 然而其线性建模能力难以捕捉复杂用户行为. 相比之下, 基于深度学习的推荐模型具备更强的表达能力, 但在联邦环境下面临更高的通信与计算开销. 为兼顾推荐性能与系统效率, 本文将所提出的差分隐私通信高效框架分别应用于上述两类主流模型, 并进行适配与优化. 第 5.1 节介绍基于子模型的联邦矩阵分解推荐算法. 第 5.2 节介绍基于差分隐私的通信高效联邦深度学习推荐算法, 其中包含浅层网络的子模型选择策略与深层网络的梯度压缩策略.

5.1 基于子模型的联邦矩阵分解推荐算法

矩阵分解推荐算法是一种广泛应用的推荐系统方法, 通过将用户-物品交互矩阵分解为用户特征矩阵和物品特征矩阵, 捕捉用户与物品之间的潜在偏好关系. 设用户-物品交互矩阵为 R , 其维度为 $m \times n$, 其中 m 表示用户数, n 表示物品数. 矩阵分解的目标是找到两个低秩矩阵 $U \in \mathbb{R}^{m \times d}$ 和 $V \in \mathbb{R}^{n \times d}$, 使得 $R \approx U \cdot V$, 其中 U 表示用户的隐含特征矩阵, V 表示物品的隐含特征矩阵, d 为隐含特征维度. 对于任意用户 i 和物品 j , 预测评分 $\tilde{r}_{ij}$ 表示为: $\tilde{r}_{ij} = u_i v_j^T$, u_i 和 v_j 分别为用户 i 和物品 j 的特征向量. 在训练过程中, 通过最小化以下目标函数来学习这些特征矩阵.

$$\min_{U, V} \sum_{(i,j) \in D} (r_{ij} - \tilde{r}_{ij}) + \lambda \|U\|_2^2 + \mu \|V\|_2^2 \quad (9)$$

其中, D 表示有评分记录的用户-物品对集合, λ 、 μ 是正则化参数, 用于防止过拟合. 在联邦矩阵分解框架下, 直接传输完整的物品特征矩阵 V 会导致通信开销过大, 并可能引发隐私问题. 为此, 本文采用 DP-SUB 子模型优化策略, 中心服务器为每个客户端选择其交互历史相关的物品子特征矩阵 V_{sub} , 并将其下发到客户端, 客户端仅更新与自身交互相关的子模型部分, 以此降低通信成本.

如图 5 所示, 中心服务器在每轮迭代中根据客户端的交互记录, 选择对应的物品子特征矩阵 V_{sub} 下发至各客户端. 客户端接收到 V_{sub} 后, 与其本地的用户特征向量 u_i 结合, 进行局部更新, 计算出新的梯度 $\tilde{g}_{\text{item}}$. 更新后的梯度经差分隐私处理后上传至中心服务器, 服务器聚合所有客户端上传的梯度信息以更新全局模型 V . 通过子模型优化策略, 客户端仅需传输相关的子模型部分, 有效减少了数据传输量, 提高了通信效率, 增强了隐私保护能力.

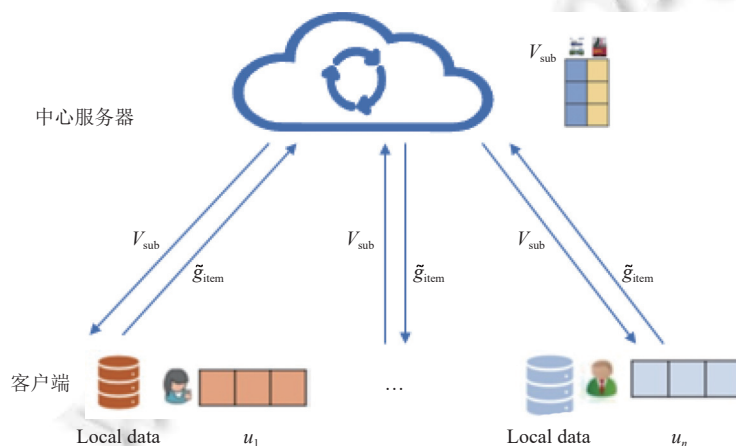


图 5 基于子模型的联邦矩阵分解推荐算法框架图

为进一步保障用户隐私, 本文在子模型优化的基础上引入梯度扰动策略. 该策略通过在客户端上传的梯度中加入拉普拉斯噪声, 以满足差分隐私需求, 防止中心服务器通过梯度信息推测用户的真实偏好. 梯度扰动步骤如下.

(1) 梯度裁剪与扰动: 为了防止单个用户梯度对全局模型产生过大影响, 首先对梯度 g_{item} 进行裁剪, 将其限制在预设阈值 C 内, 然后添加拉普拉斯噪声以满足差分隐私要求, 得到扰动后的梯度 $\tilde{g}_{\text{item}}$:

$$\tilde{g}_{\text{item}} = g_{\text{item}} \times \min\left(1, \frac{C}{\|g_{\text{item}}\|}\right) + \text{Lap}\left(\frac{\Delta f}{\epsilon_2}\right) \quad (10)$$

其中, $\|g_{\text{item}}\|$ 为梯度的 L_2 -范数, C 为裁剪阈值, Δf 为梯度的敏感度, ϵ_2 为隐私预算参数. 该公式将梯度裁剪与拉普拉斯噪声添加整合, 确保上传的梯度信息在隐私保护范围内.

(2) 梯度上传: 客户端将扰动后的梯度 $\tilde{g}_{\text{item}}$ 上传至中心服务器, 服务器在聚合过程中仅接收到已扰动的梯度信息, 实现了有效的隐私保护.

子模型优化与梯度扰动策略的结合, 使得算法在通信效率和隐私保护之间实现了有效的平衡. 通过子模型的选择与下发, 客户端仅处理与自身交互数据相关的特征部分, 极大地减少了数据传输量, 也显著降低了本地训练过程中实际参与计算的参数维度, 从而减少了总体计算开销; 梯度扰动策略则在数据传输过程中对裁剪后的子梯度加入噪声, 以增强隐私保护, 同时避免了全局梯度带来的额外计算负担. 该方法不仅在通信和隐私保护之间取得了良好权衡, 也在计算资源受限的场景中展现出较强的可扩展性.

下面概述 Priv-FedMF-Sub 算法, 该算法在 DP-SUB 子模型选择策略的基础上, 结合差分隐私机制与联邦矩阵分解技术, 通过客户端和服务器的协作训练, 实现对用户特征向量和物品特征向量的高效更新, 同时兼顾用户隐私保护与通信效率. 算法主要分为两部分: 客户端通过本地扰动用户梯度信息并上传至服务器, 服务器端对物品特征向量进行全局聚合与更新, 并将子模型分发回客户端以完成新一轮训练. 算法 2 为 Priv-FedMF-Sub (private federated matrix factorization submodel) 算法的具体步骤.

算法 2. Priv-FedMF-Sub 算法.

输入: 客户端数量 n , 采样比例 k , 全局训练轮数 T , 本地训练轮数 E , 隐私预算 ϵ_2 , 裁剪阈值 C ;

输出: 全局模型 V .

服务器端:

1. 初始化特征矩阵 V
2. for $t = 1$ to T do 全局迭代
 3. $V_{\text{sub}} \subseteq V'$ 由 DP-SUB 算法筛选下发子模型
 4. 随机抽取 m 个客户端: $m = \max(n \times k, 1)$, 其用户向量集合为 $S_t = \{S_1, S_2, \dots, S_m\}$
 5. 下发 V_{sub} 至用户向量 S_t
 6. for $i = 1$ to m do
 7. 用户向量 S_i 接收扰动梯度 $\tilde{g}_{\text{item}}$
 8. 梯度更新: $v_{\text{item}}^{t+1} = v_{\text{item}}^t - \alpha \tilde{g}_{\text{item}}$
 9. end for
 10. 服务器聚合形成新一轮全局模型 V^{t+1}
11. end for

客户端:

1. 初始化特征向量 u_i , 下载子模型 V_{sub}
 2. for $e = 1$ to E do
 3. 本地训练: $u_i^{e+1} = u_i^e - \alpha g_{u_i}^e$
 4. for each item $\in V_{\text{sub}}$ do
 5. 计算 item 梯度 g_{item}
-

6. 梯度扰动: $\tilde{g}_{\text{item}} = g_{\text{item}} \times \min\left(1, \frac{C}{\|g_{\text{item}}\|}\right) + \text{Lap}\left(\frac{\Delta f}{\epsilon_2/e}\right)$
7. 上传 $\tilde{g}_{\text{item}}$ 至服务器
8. end for
9. end for

5.2 基于差分隐私的通信高效联邦深度学习推荐算法

矩阵分解模型在联邦推荐中具有较低的计算和通信开销, 但其仅能建模用户与物品之间的线性关系, 难以有效表达更复杂的交互模式, 导致在部分实际场景下推荐性能受限. 为弥补其建模能力不足, 深度学习推荐模型因其更强的特征表达与非线性建模能力, 成为更具潜力的替代方案. 本节将在前述隐私保护与通信优化框架基础上, 进一步适配深度模型结构, 并提出结构化划分与分层优化策略, 以提升其在联邦训练中的通信效率和实用性.

5.2.1 浅层网络的子模型选择策略

在深度学习推荐系统中, 神经网络通常由多个层级组成, 根据其在网络结构中的位置和功能, 可以大致划分为浅层网络与深层网络, 其中浅层网络主要包括嵌入层以及紧随其后的 1-2 层浅层全连接层 (如第 1 层和第 2 层 MLP), 这些层在模型中承担着特征表示的关键作用. 嵌入层将高维稀疏的用户和物品数据映射到低维稠密的向量空间, 而浅层神经元层负责捕捉基础的用户偏好和物品特征. 然而, 随着用户和物品种类的增长, 嵌入层参数量随之成倍增加, 浅层神经元层的参数规模也会显著扩大. 这种情况下, 完整传输浅层网络的参数会导致通信开销过高, 尤其是在带宽受限或计费网络环境下, 这一问题尤为突出.

为了解决浅层网络通信开销过大的问题, 本文将 DP-SUB 算法扩展至基于深度学习的联邦推荐算法中, 提出了一种浅层网络的子模型选择策略, 以减少需要上传的参数数量, 从而降低通信开销. 具体而言, 该策略在每轮训练中动态选择浅层网络的关键参数与更新, 避免传输整个浅层网络的所有参数. 客户端基于本地数据更新嵌入层和浅层神经元层, 并生成局部更新 ΔX_{sub} . 中心服务器接收到子模型 ΔX_{sub} , 更新后通过 FedSubAvg 算法进行聚合和校正, 确保全局模型更新的稳定性和准确性.

如图 6 所示, 服务器接收到客户端上传的子模型更新后, 通过 FedSubAvg 算法进行聚合和校正. FedSubAvg 算法通过对参与客户端上传参数的加权平均实现模型参数的整合, 具体公式如下:

$$E_C[\Delta X] = \frac{1}{n_m} \sum_{i=1}^N \Delta X_{i,m} \quad (11)$$

其中, n_m 表示参与更新的客户端数量. 通过引入系数校正这种精细的调整, FedSubAvg 算法进一步保证了模型更新的准确性和效率, 使得浅层网络的通用信息能够得到更加高效的学习和利用.

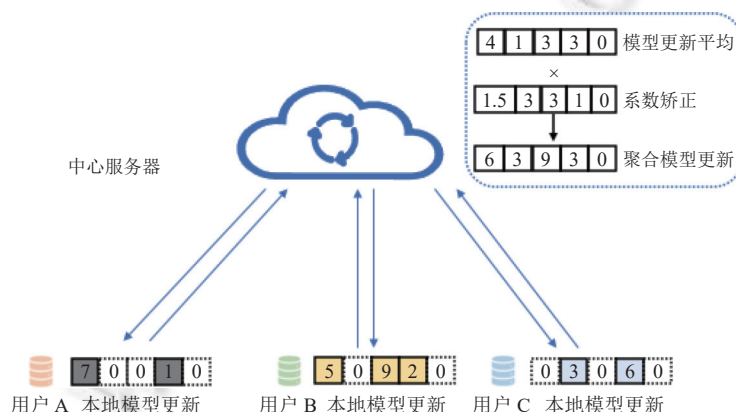


图 6 基于 FedSubAvg 的子模型聚合优化

5.2.2 深层网络的梯度压缩策略

在深度学习推荐系统中, 深层网络通常由多层非线性全连接层(如第 3 层及之后的 MLP 层)构成, 其主要功能是进一步建模用户和物品之间的高阶非线性交互关系. 深层网络的梯度矩阵因参数规模庞大而成为通信开销的主要来源. 深度学习模型往往需要传输数百万至数十亿的参数梯度, 尤其在推荐系统中, 用户和物品数量极为庞大, 嵌入层和神经元层的参数随规模成倍增长, 导致梯度矩阵的数据量进一步膨胀. 在联邦学习场景中, 客户端每轮需上传本地训练后的完整梯度至服务器进行全局聚合, 这一过程伴随的大量梯度数据传输, 不仅显著增加了通信延迟, 还对带宽和存储提出了严苛要求. 频繁的梯度更新进一步加剧了通信成本问题, 影响了模型的收敛速度与整体性能. 因此, 设计高效的梯度压缩策略以缓解通信开销成为联邦学习推荐系统的关键挑战.

为解决上述问题, 本文提出一种梯度压缩策略, 从优化梯度传输角度全面降低通信开销. 在深度学习模型中, 梯度矩阵往往表现出长尾分布的特性, 即仅少数梯度元素对模型性能优化起到关键作用, 而大部分梯度幅值较小, 其在参数更新中的贡献有限, 甚至可能引入额外的噪声干扰^[46]. 若直接传输完整的梯度矩阵不仅浪费了大量的通信资源, 还会导致无效梯度信息干扰模型优化过程. 为此, 本文设计了一种两阶段的梯度压缩机制, 分别包括稀疏化选择与梯度裁剪与扰动. 该机制一方面减少冗余梯度的传输量, 降低通信负担; 另一方面通过控制梯度幅值范围并注入噪声, 有效实现差分隐私保护. 在稀疏化过程中, 通过动态计算阈值 τ , 筛选出幅值绝对值较大的关键梯度, 其余元素置零. 具体定义如下:

$$\nabla w'_{ij} = \begin{cases} \nabla w_{ij}, & |\nabla w_{ij}| \geq \tau \\ 0, & |\nabla w_{ij}| < \tau \end{cases} \quad (12)$$

其中, 阈值 τ 可以通过梯度幅值分布的动态统计特性进行调整. 例如, 选择梯度绝对值排名前 k 的元素, 或者通过设定稀疏比例 $\text{sparsity} = k/n$ 的方式确定, 保留对模型优化贡献最大的梯度. 稀疏化有效减少了传输梯度的数量, 在降低通信成本的同时保留了模型优化的关键信息. 为了进一步控制梯度幅值范围并提升通信稳定性, 稀疏化后的梯度矩阵需进行裁剪操作. 梯度裁剪的目标是限制过大梯度对模型训练和传输效率的负面影响. 裁剪操作通过约束梯度矩阵的 L_2 范数实现, 数学表达为:

$$\nabla W_{\text{clip}} = \nabla W' \times \min\left(1, \frac{C}{\|\nabla W'\|_2}\right) \quad (13)$$

当梯度的 L_2 -范数超过预设值 C 时, 裁剪操作对梯度进行缩放, 使其幅值不会超过预设的上限. 此过程不仅减少了过大梯度引起的不稳定性, 还为后续的隐私保护提供了计算基础. 其次, 为保障隐私安全并防止通过梯度逆推出用户数据, 本文对 ∇W_{clip} 进一步进行裁剪并添加高斯噪声实现差分隐私保护. 具体数学表达为:

$$\widetilde{\nabla W} = \nabla W_{\text{clip}} \times \min\left(1, \frac{P}{\|\nabla W_{\text{clip}}\|_2}\right) + N(0, \sigma^2 P^2) \quad (14)$$

其中, 高斯噪声 $N(0, \sigma^2 P^2)$ 为零均值、方差为 $\sigma^2 P^2$ 的正态分布噪声, σ 为噪声幅度参数. 确保每轮更新中的梯度满足 ϵ -差分隐私保护要求. 即使攻击者获得上传的梯度信息, 也难以准确重构用户的敏感数据.

通过稀疏化、裁剪及高斯噪声保护的结合, 本文的梯度压缩策略显著减少了传输梯度的体积和通信成本, 有效缓解了联邦学习中通信资源不足的问题, 同时保证了梯度传输的稳定性与隐私保护. 在计算开销方面, 客户端每轮仅更新子模型对应的浅层网络参数与稀疏梯度, 显著减少了局部计算负担, 避免了深层全模型参与所带来的资源开销. 该策略在通信效率、隐私保护与计算复杂度之间实现了良好平衡, 为联邦学习推荐系统在大规模场景下的高效部署提供了理论支持.

5.2.3 Priv-FedNCF-Sub-Compress 算法概述

该算法通过服务器端与客户端的协作, 在保障用户隐私的同时降低通信成本. 具体而言, 服务器端负责选择适配的子模型与深层网络模型参数, 并分发至客户端进行本地训练; 客户端在训练过程中对梯度和模型更新进行裁剪并加入差分隐私噪声后, 将扰动后的模型参数上传至服务器端进行聚合更新. 算法引入了 DP-SUB 子模型选择

策略, 用于动态选择网络结构中的深层子模型以优化通信开销, 同时采用本地梯度裁剪与高斯噪声机制进一步增强隐私保护能力. 算法 3 为 Priv-FedNCF-Sub-Compress (privacy compressed submodel of federated neural collaborative filtering) 算法的具体步骤.

算法 3. 基于差分隐私的通信高效联邦深度学习推荐算法 Priv-FedNCF-Sub-Compress.

输入: 客户端数量 N , 采样比例 k , 全局训练轮数 T , 本地训练轮数 E , 参与更新的客户端数量 n_m , 隐私预算 ϵ_3 , 本地裁剪阈值 C , 稀疏化阈值 τ , 梯度扰动裁剪阈值 P ;

输出: 全局模型 W .

服务器端:

1. 初始化全局参数 $W_0 = W_s + W_d$
2. for $t = 1$ to T do 全局迭代
3. 由 DP-SUB 算法筛选下发浅层网络 W_s 的子模型 X_{sub}
4. 随机抽取 m 个客户端: $m = \max(n \times k, 1)$, 其用户向量集合为 $S_t = \{S_1, S_2, \dots, S_m\}$
5. 下发 X_{sub} 和深层网络参数 W_d 至用户向量 S_t
6. for $i = 1$ to m do
7. 用户向量 S_i 接收扰动更新参数 $\nabla X'_{\text{sub}}$ 和 $\nabla W'_{\text{clip}}$
8. end for
9. 服务器采用 FedSubAvg 形成新一轮的全局模型 W^{t+1}
10. end for

客户端:

1. 初始化特征向量 u_i , 下载子模型 X_{sub} 和 W_d
 2. for $e = 1$ to E do
 3. 本地训练: $u_i^{e+1} = u_i^e - \alpha g_{u_i}^e$ 计算 ∇X_{sub} 和 ∇W_d
 4. 根据公式 (13) 计算深层网络梯度 ∇W_{clip}
 5. 计算扰动后的 $\nabla X'_{\text{sub}} = \nabla X_{\text{sub}} \times \min\left(1, \frac{P}{\|\nabla X_{\text{sub}}\|_2}\right) + N(0, \sigma^2 P^2)$ 进行本地裁剪并加噪
 6. 计算扰动后的 $\nabla W'_{\text{clip}} = \nabla W_{\text{clip}} \times \min\left(1, \frac{P}{\|\nabla W_{\text{clip}}\|_2}\right) + N(0, \sigma^2 P^2)$ 进行本地裁剪并加噪
 7. 上传 $\nabla X'_{\text{sub}}$ 和 $\nabla W'_{\text{clip}}$ 至服务器
 8. end for
-

6 实验分析

本节对提出方法的性能进行了系统验证与全面评估. 首先, 第 6.1 节介绍实验中使用的数据集, 包括其来源、规模及主要特性, 为实验研究提供了基础支持. 第 6.2 节详细说明所采用的评价指标及实验中选取的基准模型, 为后续性能对比分析奠定了标准. 第 6.3 节描述实验的具体配置, 包括模型参数、训练过程及实验硬件环境等技术细节. 第 6.4 节基于实验结果对方法性能进行深入分析, 并与基准模型进行对比, 验证所提方法的有效性与优越性.

6.1 实验数据集

本文使用 MovieLens 100K 和 MovieLens 1M 两个广泛应用于推荐系统研究的标准化数据集. 这两个数据集分别代表了小规模与大规模推荐场景, 适用于不同条件下评估算法在处理稀疏数据与高维交互任务中的表现. 表 1 给出了数据集所对应的详细信息.

表 1 实验数据集

数据集	MovieLens 100K	MovieLens 1M
用户数量	943	6040
电影数量	1682	3900
评分数量	100 000	1 000 209
用户平均评分数	106.04	165.57
电影平均被评分数	59.45	256.47
评分矩阵密度 (%)	6.30	4.25

MovieLens 100K 数据集包含 943 名用户对 1682 部电影的 100 000 条评分记录, 评分范围为 1-5 的整数, 评分矩阵的稀疏度为 6.30%。此外, 数据集还包括用户的基本属性 (如性别、年龄、职业和邮编) 以及电影的分类标签 (如流派)。该数据集的特点是规模较小, 稀疏性较高, 适合快速验证算法在稀疏数据场景中的表现及其在小规模场景中的收敛性。

MovieLens 1M 数据集包含 6040 名用户对 3900 部电影的 1 000 209 条评分记录, 评分矩阵的稀疏度为 4.25%, 相较于 MovieLens 100K 显著降低。用户属性更为详细, 除性别、年龄、职业外, 还涵盖了更多样化的电影标签信息。该数据集的规模更大、评分密度更高, 能够更有效地评估算法在复杂和高维场景中的扩展能力和鲁棒性。

在实验中, 两个数据集按照用户进行划分, 每个用户的数据被视为一个独立的客户端。客户端仅保留本地评分记录及相关属性数据, 并按照 8:2 的比例将数据分为训练集和测试集, 以模拟分布式环境下的数据训练与测试过程。

6.2 评价指标及基准模型

在本文中, 我们采用 HR@10 和 NDCG@10 作为衡量推荐算法可用性的核心评价指标, 并引入传输总参数量作为评估通信开销的重要指标, 从推荐性能与通信效率两个维度对算法进行全面评估。

HR@10 (hit ratio at 10) 是用来衡量推荐算法覆盖用户实际兴趣的能力的重要指标, 其定义为推荐列表中是否包含用户真实喜好的条目。该指标能够直观反映推荐算法的实用性和有效性, 值越高说明推荐结果越符合用户需求。NDCG@10 (normalized discounted cumulative gain at 10) 则进一步考虑了推荐条目的排名位置对用户体验的影响, 其值越高表明推荐算法对用户兴趣的排序更精准, 从而提升了推荐结果的用户满意度。传输总参数量反映了联邦学习过程中客户端与服务器之间交换的参数数据总量, 参数量越小表明通信效率越高, 有助于减少带宽占用和降低通信成本。

实验针对所提出的两种方法, 分别选取了目前具有代表性的基线方法进行对比。在矩阵分解推荐方法中, 基线方法包括: (1) FedMF^[23]: 一种传统的联邦矩阵分解推荐算法, 所有客户端按照传统联邦学习的方式协同训练全局模型; (2) DP-FedMF^[17]: 一种具有本地差分隐私保护的联邦矩阵分解推荐算法, 客户端在本地上传梯度时加入满足差分隐私的噪声。在深度学习推荐方法中, 基线方法包括: (1) FedNCF^[23]: 一种传统的联邦深度学习推荐算法, 所有客户端按照传统联邦学习的方式协同训练全局模型; (2) DP-FedNCF^[55]: 一种具有本地差分隐私保护的联邦深度学习推荐算法, 客户端在本地上传梯度时加入满足差分隐私的高斯噪声; (3) DP-FedNCF-DGC^[38]: 在 DP-FedNCF 的基础上采用深度梯度压缩策略, 以显著降低通信成本并保持模型性能。上述基线方法覆盖了联邦矩阵分解和深度学习推荐场景下的主流方法, 为实验提供了多角度的对比参考, 验证了所提方法在隐私保护和通信效率方面的综合优势。

6.3 实现细节

实验环境基于 Ubuntu 20.04.6 LTS 操作系统, 硬件配置包括 2 块 32 GB V100 GPU, Intel(R) Xeon(R) Gold 5118 CPU@2.30 GHz 处理器, 内存容量为 256 GB。实验中, 潜在特征向量维度均设置为 32, 批量大小分别为 256 (MovieLens 100K) 和 512 (MovieLens 1M), 负样本比例为 1:4。每轮通信时, 分别从 MovieLens 100K 和 MovieLens 1M 中采样 100 和 200 个客户端参与本地训练, 每个客户端进行 5 轮本地迭代后进行全局模型更新, 通信总轮数为 400。优化器方面, 矩阵分解算法采用 SGD 优化器, 学习率设置为 0.001, 深度学习算法采用 Adam 优化器, 学习率设置为 0.0005。梯度裁剪阈值统一设置为 $C=1.0$, 客户端隐私预算设置为 2。本实验所有程序代码使用 Python 语言编写 (相关的实验代码公开发布于 https://github.com/Drxdx/FL_recommend)。

6.4 实验结果与分析

本节在 2 个数据集上对 Priv-FedMF-Sub 和 Priv-FedNCF-Sub-Compress 进行全面评估, 并与前述基准方法对比。

● Priv-FedMF-Sub 有效性评估

如图 7 和图 8 分别展示了 3 种算法在 MovieLens 100K 和 MovieLens 1M 数据集上的推荐准确性。通过 HR@10 和 NDCG@10 两个指标结果可以看出, FedMF 在这两个指标上的表现最为优异。该算法通过不直接上传用户原始数据的方式, 利用传统的联邦学习框架协同训练全局模型, 从而有效避免了用户数据的直接泄露。然而, 尽管 FedMF 在推荐准确性上表现突出, 其在联邦训练过程中通过交换中间梯度的方式仍然存在隐私泄露的潜在风险, 无法确保绝对的隐私安全性。为增强隐私保护, DP-FedMF 在 FedMF 的基础上引入了梯度扰动机制, 满足差分隐私的要求。尽管该方法提高了隐私保护性, 然而由于噪声的加入, DP-FedMF 在推荐准确性上有所下降。与 DP-FedMF 相比, Priv-FedMF-Sub 在保持相同隐私保护程度的前提下, 推荐准确性未出现显著下降, 表现出与 DP-FedMF 相当的模型可用性。

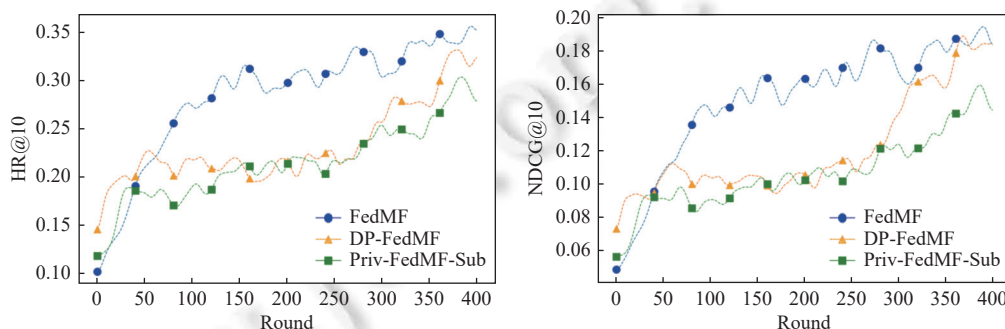


图 7 MovieLens 100K 数据集下推荐算法模型性能

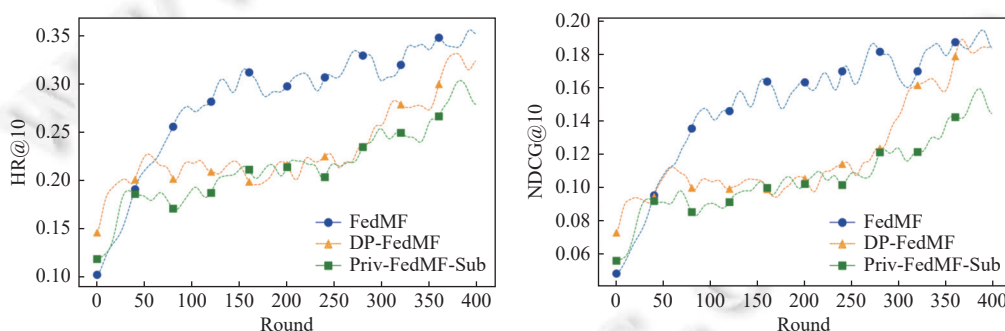


图 8 MovieLens 1M 数据集下推荐算法模型性能

本文重点关注算法的通信开销, 图 9 展示了 FedMF、DP-FedMF 和 Priv-FedMF-Sub 这三种算法的通信效率对比。在传统联邦学习框架下, FedMF 和 DP-FedMF 均采用完整的全局模型进行参数传输, 因此其传输的总参数量相同, 其中使用整个训练过程中所需的总参数传输量来衡量通信成本, 单位为百万 (M)。相比之下, Priv-FedMF-Sub 通过采用子模型替代全局模型, 显著降低了通信开销。具体而言, 在 MovieLens 100K 数据集上, FedMF 的总参数传输量为 2154.24M, 而 Priv-FedMF-Sub 为 698.67M, 通信效率提升 67.57%; 在 MovieLens 1M 数据集上, FedMF 的总传输参数量为 9487.36M, Priv-FedMF-Sub 为 2907.846M, 通信效率提升 69.35%。这种差异源于 Priv-FedMF-Sub 通过选择适当的子模型, 依据交互频率均值设置阈值, 从而减少了需要传输的参数数量。值得注意的是, 由于 MovieLens 1M 数据集中的电影数量更多, 相应全局模型 (特征向量矩阵) 也更大, 因此其通信开销明显高于 MovieLens 100K 数据集。这一结果表明, Priv-FedMF-Sub 方法在保证隐私保护的同时, 能够有效降低通信成本。

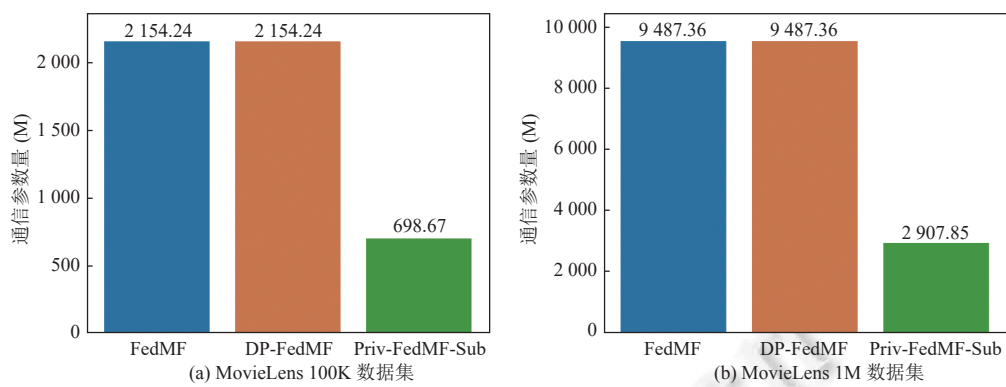


图 9 通信开销对比图

此外, 为了进一步评估本文方法的通信效率, 实验中我们还将其与几种常见的矩阵分解推荐算法中的通信开销降低策略进行了对比, 实验结果如表 2 所示. 具体而言, DP-FedMF 采用全局模型进行参数传输, 隐私保护水平与本文所提出的方法相当. DP-FedMF-SVD 和 DP-FedMF-TopK 分别在 DP-FedMF 的基础上引入了 SVD 和 Top-K 压缩策略, 而 Priv-FedMF-Sub 则为本文提出的改进算法. 由表 2 可见, 尽管与 DP-FedMF 相比, 其他方法在推荐准确性上均存在不同程度的下降 (这主要源于压缩操作或仅训练部分子模型导致的精度损失), 然而这些方法在通信效率上均有显著提升. 值得注意的是, 虽然 Priv-FedMF-Sub 方法在 HR@10 指标上略有下降, 但其在 NDCG@10 指标上有所提升, 表明本文方法在排序较高的推荐项上表现出更高的精确性. 同时, Priv-FedMF-Sub 在通信效率上的提升优于其他压缩策略, 进一步验证了其在减少通信开销方面的优势.

表 2 性能对比图

方法	HR@10	NDCG@10	总参数传输量 (M)
FedMF	0.650	0.367	9487.36
DP-FedMF	0.515	0.293	9487.36
DP-FedMF-SVD	0.452	0.248	4785.93
DP-FedMF-TopK	0.453	0.250	4943.68
Priv-FedMF-Sub	0.435	0.274	2907.85

● Priv-FedMF-Sub 消融实验

为深入分析 Priv-FedMF-Sub 算法中各组成部分对整体性能的贡献, 本文进行了消融实验, 评估了各关键模块的效果. 如图 10 和图 11 所示, 我们比较了 3 种不同配置的算法: 1) FedMF-Sub-RR: 去除基于最大似然估计的真实频率估计和梯度扰动策略; 2) FedMF-Sub-RR-EST: 仅去除梯度扰动策略; 3) Priv-FedMF-Sub: 即本文提出的完整算法.

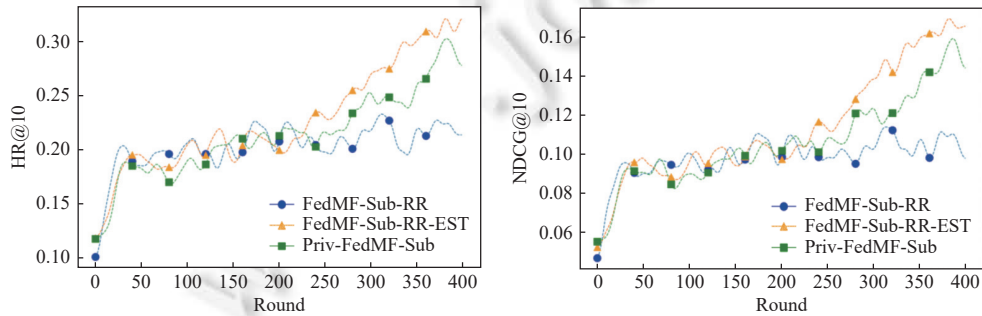


图 10 MovieLens 100K 数据集下消融实验对比图

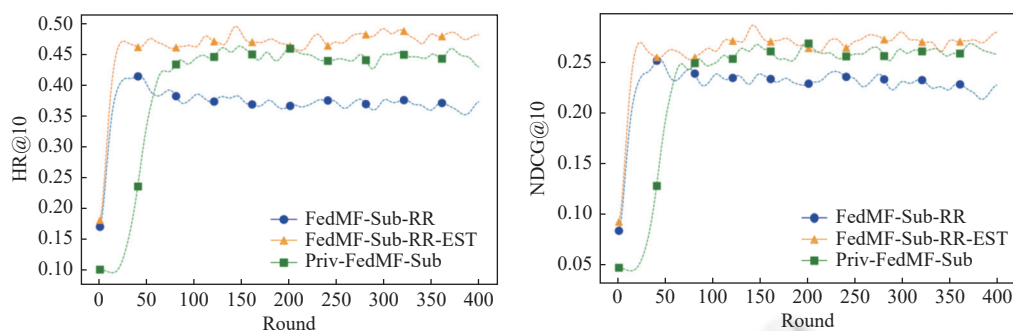


图 11 MovieLens 1M 数据集下消融实验对比图

通过对 $HR@10$ 和 $NDCG@10$ 两个指标进行评估,可以看出,与 FedMF-Sub-RR 相比, FedMF-Sub-RR-EST 在模型可用性上显著提升,这表明 FedMF-Sub-RR 由于直接基于扰动数据进行子模型选择,导致模型偏差增大,降低了可用性。而在 FedMF-Sub-RR-EST 中,服务器端通过最大似然估计补偿了扰动数据的影响,从而有效提升了推荐质量。在此基础上, Priv-FedMF-Sub 虽然在推荐准确性上有所下降,但其通过引入梯度扰动策略,增强了对用户隐私的保护,显著降低了通过梯度反推用户数据的隐私泄露风险。

综上所述,消融实验结果表明,基于差分隐私的子模型选择和梯度扰动策略在提升推荐准确性和保护用户隐私方面发挥了重要作用。其中,最大似然估计方法在子模型选择中的应用显著提升了模型可用性,而梯度扰动策略则有效加强了隐私保护,二者结合不仅提高了通信效率,也保障了隐私安全。

• Priv-FedNCF-Sub-Compress 有效性评估

图 12 和图 13 分别展示了 4 种算法在 MovieLens 100K 和 MovieLens 1M 数据集上的推荐效果。结果表明, FedNCF 模型在未采取隐私保护措施的情况下,推荐准确性表现最佳。然而,尽管联邦学习通过避免集中存储用户数据减少了隐私泄露的风险,客户端上传的中间梯度仍可能泄露用户隐私。引入差分隐私保护的 DP-FedNCF 算法在 $HR@10$ 和 $NDCG@10$ 指标上有所下降,表明梯度噪声的引入对模型准确性产生了负面影响。相比之下, DP-FedNCF-DGC 在进一步减少通信开销的过程中引入了深度梯度压缩策略,导致推荐准确性进一步降低,因为部分有效信息在压缩过程中丢失。与 DP-FedNCF 不同,本文提出的 Priv-FedNCF-Sub-Compress 算法在保持相同隐私保护水平的情况下,推荐准确性未发生显著下降,且与 DP-FedNCF 相当。同时, Priv-FedNCF-Sub-Compress 较 DP-FedNCF-DGC 在推荐准确性上更具优势。这是因为,浅层网络通常包含更为重要的通用信息,而 DP-FedNCF-DGC 对整个网络进行了压缩,而 Priv-FedNCF-Sub-Compress 仅对深层网络进行梯度压缩,从而有效保留了浅层网络的关键信息。

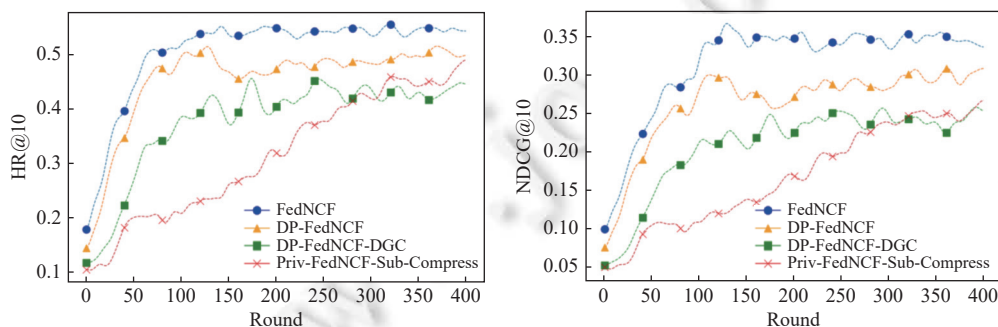


图 12 MovieLens 100K 数据集下推荐算法模型性能

在通信开销方面,图 14 对 FedNCF、DP-FedNCF、DP-FedNCF-DGC 及 Priv-FedNCF-Sub-Compress 这 4 种算法的通信成本进行比较。FedNCF 与 DP-FedNCF 均采用传统联邦训练框架下的完整全局模型进行参数传输,因

此在每轮迭代中的参数传输量相同. 与之相比, Priv-FedNCF-Sub-Compress 在两个数据集上的通信效率均显著提高, 具体表现为平均传输参数数量的减少 (单位为 M). 在 MovieLens 100K 数据集上, FedNCF 每轮迭代的平均传输量为 10 771.2M, 而 Priv-FedNCF-Sub-Compress 则降至 4 200.0M, 实现了约 61% 的通信效率提升; 在 MovieLens 1M 数据集上, FedNCF 的传输量为 47 436.8M, 而 Priv-FedNCF-Sub-Compress 降至 4 421.2M, 实现了约 90.7% 的通信效率提升.

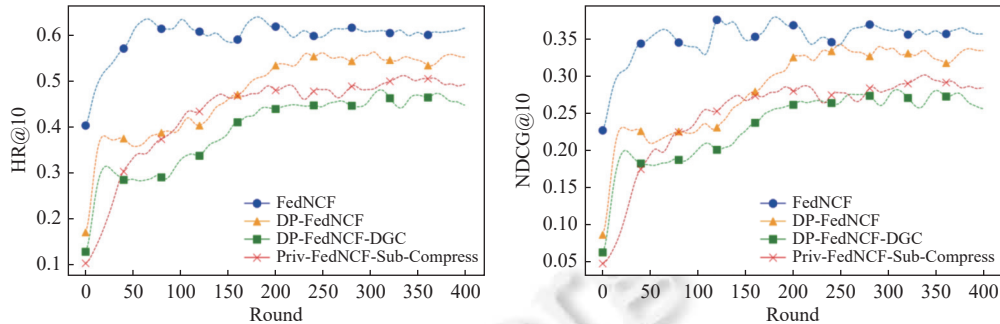


图 13 MovieLens 1M 数据集下推荐算法模型性能

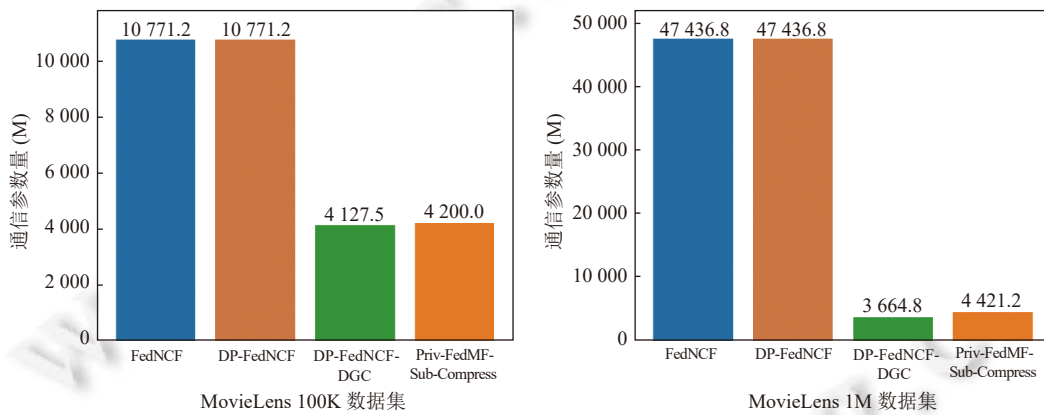


图 14 通信开销对比图

这一显著优化源于 Priv-FedNCF-Sub-Compress 算法中结合了浅层网络的子模型选择策略与深层网络的梯度压缩策略, 具体通过对物品特征嵌入层进行子模型筛选以及对深层神经网络的梯度进行压缩, 从而有效减轻了联邦训练中的通信负担. 与采用深度梯度压缩策略的 DP-FedNCF-DGC 相比, Priv-FedNCF-Sub-Compress 在达成相似通信效率的同时, 推荐准确性更高, 进一步验证了本方法的有效性与优越性.

● Priv-FedNCF-Sub-Compress 消融实验

为深入分析 Priv-FedNCF-Sub-Compress 算法中各组成部分对整体性能的贡献, 本文进行了消融实验, 评估了不同算法配置的效果. 实验中, 我们对比了 3 种算法配置: 1) FedNCF: 传统联邦深度学习推荐算法; 2) Priv-FedNCF-Sub: 去除深层网络压缩策略, 仅采用浅层网络子模型策略; 3) Priv-FedNCF-Sub-Compress: 结合浅层网络子模型选择与深层网络梯度压缩策略的完整方法.

图 15 和图 16 展示了在 MovieLens 100K 和 MovieLens 1M 数据集上, 基于 HR@10 和 NDCG@10 指标的结果. 相较于基线算法 FedNCF, Priv-FedNCF-Sub 在推荐准确性上有所下降, 原因在于其通过扰动交互数据与梯度来保护隐私, 同时应用子模型策略以降低通信开销, 虽牺牲部分准确性, 但有效提升了隐私保护与通信效率. 相比之下, Priv-FedNCF-Sub-Compress 进一步引入了深层网络的梯度压缩策略, 虽然导致准确性有所下降, 但在显著降

低通信开销的同时, 仍保持了较为优良的推荐性能, 显示了其在隐私保护与通信效率上的优越性。

综上所述, 消融实验结果验证了 Priv-FedNCF-Sub-Compress 算法各部分的关键作用. 通过对深度神经网络进行结构划分, 分别在浅层和深层网络中应用子模型策略与梯度压缩策略, 有效降低了联邦推荐中的通信负担. 结合基于差分隐私的子模型选择和梯度扰动策略, 进一步增强了用户数据的隐私保护. 该算法在保证隐私安全的同时, 显著提高了通信效率, 并未显著影响模型可用性。

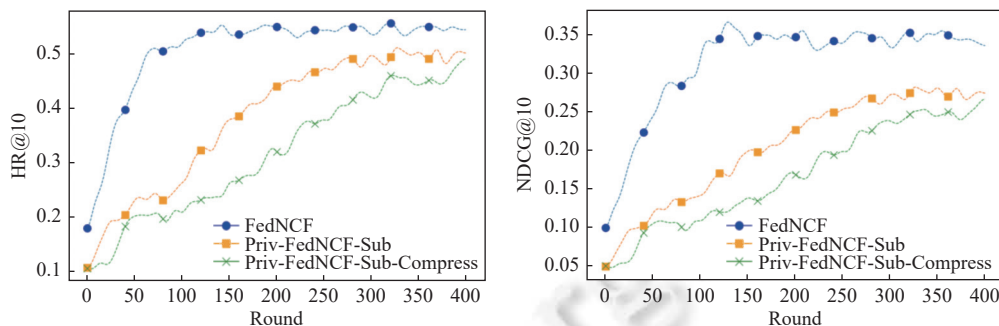


图 15 MovieLens 100K 数据集下消融实验对比图

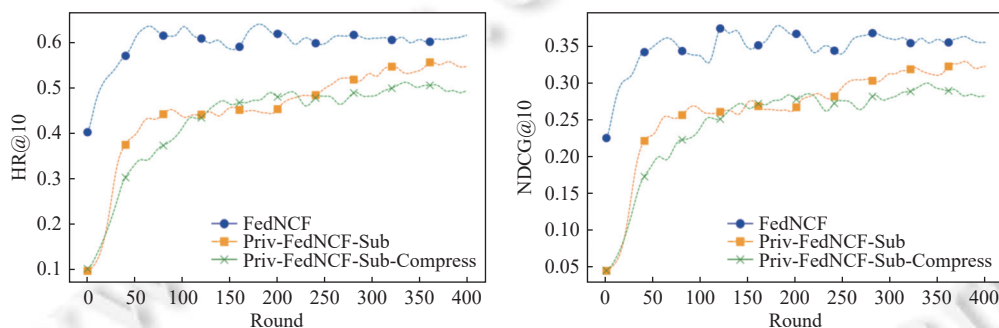


图 16 MovieLens 100K 数据集下消融实验对比图

7 总 结

本文针对传统联邦推荐系统在通信开销和隐私保护方面的挑战, 提出了两种基于差分隐私的通信高效联邦推荐算法: 基于矩阵分解的 Priv-FedMF-Sub 算法和基于深度学习的 Priv-FedNCF-Sub-Compress 算法. 首先, 通过引入子模型策略和最大似然估计方法优化, 平衡了隐私保护与模型效用之间的关系, 确保了推荐系统的高效性. 其次, 在联邦深度学习推荐中, 通过结构化划分全局模型并为浅层和深层网络定制优化策略, 包括梯度压缩和基于差分隐私的梯度扰动技术, 有效减轻了通信负担. 最后, 从理论和实验两方面验证了本文方案的安全性、有效性和鲁棒性. 由于通信效率和隐私保护之间的平衡仍是一个挑战, 后续研究可以考虑通过结合更精细的压缩技术 (如参数剪枝、量化、模型蒸馏等) 来进一步提升通信效率. 同时, 本文拟在未来工作中对更先进的隐私保护技术, 如同态加密、安全多方计算等进行探讨, 进一步强化数据保护能力并提升系统的安全性。

References

- [1] Shi Y, Larson M, Hanjalic A. Collaborative filtering beyond the user-item matrix: A survey of the state of the art and future challenges. ACM Computing Surveys (CSUR), 2014, 47(1): 3. [doi: 10.1145/2556270]
- [2] Zhang RZ, Xie XH, Mao JX, Liu YQ, Zhang M, Ma SP. Constructing a comparison-based click model for Web search. In: Proc. of the 2021 Web Conf. Ljubljana: ACM, 2021. 270–283. [doi: 10.1145/3442381.3449918]
- [3] Pan K, Ong YS, Gong MG, Li H, Qin AK, Gao Y. Differential privacy in deep learning: A literature survey. Neurocomputing, 2024, 589:

127663. [doi: [10.1016/j.neucom.2024.127663](https://doi.org/10.1016/j.neucom.2024.127663)]
- [4] Narayanan A, Shmatikov V. Robust de-anonymization of large sparse datasets. In: Proc. of the 2008 IEEE Symp. on Security and Privacy. Oakland: IEEE, 2008. 111–125. [doi: [10.1109/SP.2008.33](https://doi.org/10.1109/SP.2008.33)]
 - [5] Yang L, Tan B, Zheng VW, Chen K, Yang Q. Federated recommendation systems. In: Yang Q, Fan LX, Yu H, eds. Federated Learning: Privacy and Incentive. Cham: Springer, 2020. 225–239. [doi: [10.1007/978-3-030-63076-8_16](https://doi.org/10.1007/978-3-030-63076-8_16)]
 - [6] Guo L, Lu ZA, Yu JL, Nguyen QVH, Yin HZ. Prompt-enhanced federated content representation learning for cross-domain recommendation. In: Proc. of the 2024 ACM Web Conf. Singapore: ACM, 2024. 3139–3149. [doi: [10.1145/3589334.3645337](https://doi.org/10.1145/3589334.3645337)]
 - [7] Lin ZH, Huang W, Zhang HY, Xu JY, Liu WM, Liao XT, Wang F, Wang SP, Tan YC. Enhancing dual-target cross-domain recommendation with federated privacy-preserving learning. In: Proc. of the 33rd Int'l Joint Conf. on Artificial Intelligence. Jeju, 2024. 238. [doi: [10.24963/ijcai.2024/238](https://doi.org/10.24963/ijcai.2024/238)]
 - [8] Muhammad K, Wang QQ, O'Reilly-Morgan D, Tragos E, Smyth B, Hurley N, Geraci J, Lawlor A. FedFast: Going beyond average for faster training of federated recommender systems. In: Proc. of the 26th ACM SIGKDD Int'l Conf. on Knowledge Discovery & Data Mining. ACM, 2020. 1234–1242. [doi: [10.1145/3394486.3403176](https://doi.org/10.1145/3394486.3403176)]
 - [9] McMahan B, Moore E, Ramage D, Hampson S, Agüera y Arcas B. Communication-efficient learning of deep networks from decentralized data. In: Proc. of the 20th Int'l Conf. on Artificial Intelligence and Statistics. Fort Lauderdale: PMLR, 2017. 1273–1282.
 - [10] Hamer J, Mohri M, Suresh A T. FedBoost: A communication-efficient algorithm for federated learning. In: Proc. of the 37th Int'l Conf. on Machine Learning. PMLR, 2020. 3973–3983.
 - [11] Niu CY, Wu F, Tang SJ, Hua LF, Jia RF, Lv CF, Wu ZH, Chen GH. Billion-scale federated learning on mobile clients: A submodel design with tunable privacy. In: Proc. of the 26th Annual Int'l Conf. on Mobile Computing and Networking. London: ACM, 2020. 31. [doi: [10.1145/3372224.3419188](https://doi.org/10.1145/3372224.3419188)]
 - [12] Chai D, Wang LY, Chen K, Yang Q. Secure federated matrix factorization. IEEE Intelligent Systems, 2021, 36(5): 11–20. [doi: [10.1109/MIS.2020.3014880](https://doi.org/10.1109/MIS.2020.3014880)]
 - [13] Song CZ, Ristenpart T, Shmatikov V. Machine learning models that remember too much. In: Proc. of the 2017 ACM SIGSAC Conf. on Computer and Communications Security. Dallas: ACM, 2017. 587–601. [doi: [10.1145/3133956.3134077](https://doi.org/10.1145/3133956.3134077)]
 - [14] Zhang SJ, Yuan W, Yin HZ. Comprehensive privacy analysis on federated recommender system against attribute inference attacks. IEEE Trans. on Knowledge and Data Engineering, 2024, 36(3): 987–999. [doi: [10.1109/TKDE.2023.3295601](https://doi.org/10.1109/TKDE.2023.3295601)]
 - [15] Truex S, Liu L, Chow KH, Gursoy ME, Wei WQ. LDP-Fed: Federated learning with local differential privacy. In: Proc. of the 3rd ACM Int'l Workshop on Edge Systems, Analytics and Networking. Heraklion: ACM, 2020. 61–66. [doi: [10.1145/3378679.3394533](https://doi.org/10.1145/3378679.3394533)]
 - [16] Buyukates B, Ulukus S. Timely communication in federated learning. In: Proc. of the 2021 IEEE Conf. on Computer Communications Workshops (INFOCOM WKSHPS). Vancouver: IEEE, 2021: 1–6. [doi: [10.1109/INFOCOMWKSHPS51825.2021.9484497](https://doi.org/10.1109/INFOCOMWKSHPS51825.2021.9484497)]
 - [17] Li ZT, Ding BL, Zhang C, Li N, Zhou J. Federated matrix factorization with privacy guarantee. Proc. of the VLDB Endowment, 2021, 15(4): 900–913. [doi: [10.14778/3503585.3503598](https://doi.org/10.14778/3503585.3503598)]
 - [18] Rong DZ, He QM, Chen JH. Poisoning deep learning based recommender model in federated learning scenarios. In: Proc. of the 31st Int'l Joint Conf. on Artificial Intelligence (IJCAI 2022). IJCAI, 2022. 2204–2210.
 - [19] Wang QY, Yin HZ, Chen T, Yu JL, Zhou A, Zhang XL. Fast-adapting and privacy-preserving federated recommender system. The VLDB Journal, 2022, 31(5): 877–896. [doi: [10.1007/s00778-021-00700-6](https://doi.org/10.1007/s00778-021-00700-6)]
 - [20] Fung BCM, Wang K, Chen R, Yu PS. Privacy-preserving data publishing: A survey of recent developments. ACM Computing Surveys (CSUR), 2010, 42(4): 14. [doi: [10.1145/1749603.1749605](https://doi.org/10.1145/1749603.1749605)]
 - [21] Salakhutdinov R, Mnih A. Probabilistic matrix factorization. In: Proc. of the 21st Int'l Conf. on Neural Information Processing Systems. Vancouver: Curran Associates Inc., 2007. 1257–1264.
 - [22] Yin HZ, Wang QY, Zheng K, Li ZX, Yang JL, Zhou XF. Social influence-based group representation learning for group recommendation. In: Proc. of the 35th IEEE Int'l Conf. on Data Engineering (ICDE). Macao: IEEE, 2019: 566–577. [doi: [10.1109/ICDE.2019.00057](https://doi.org/10.1109/ICDE.2019.00057)]
 - [23] Perifanis V, Efraimidis PS. Federated neural collaborative filtering. Knowledge-based Systems, 2022, 242: 108441. [doi: [10.1016/j.knosys.2022.108441](https://doi.org/10.1016/j.knosys.2022.108441)]
 - [24] Lin GY, Liang F, Pan WK, Ming Z. FedRec: Federated recommendation with explicit feedback. IEEE Intelligent Systems, 2021, 36(5): 21–30. [doi: [10.1109/MIS.2020.3017205](https://doi.org/10.1109/MIS.2020.3017205)]
 - [25] Duan SJ, Zhang DY, Wang YB, Li LX, Zhang YX. JointRec: A deep-learning-based joint cloud video recommendation framework for mobile IoT. IEEE Internet of Things Journal, 2020, 7(3): 1655–1666. [doi: [10.1109/JIOT.2019.2944889](https://doi.org/10.1109/JIOT.2019.2944889)]
 - [26] Huang W, Liu J, Li TR, Huang TQ, Ji SG, Wan JH. FedDSR: Daily schedule recommendation in a federated deep reinforcement learning

- framework. *IEEE Trans. on Knowledge and Data Engineering*, 2023, 35(4): 3912–3924. [doi: [10.1109/TKDE.2021.3130265](https://doi.org/10.1109/TKDE.2021.3130265)]
- [27] Wang H, Zhou H, Lit MZ, Zhao L, Leung VCM. Federated distributed deep reinforcement learning for recommendation-enabled edge caching. In: *Proc. of the 2024 IEEE Conf. on Computer Communications Workshops (INFOCOM WKSHPS)*. Vancouver: IEEE, 2024. 1–2. [doi: [10.1109/INFOCOMWKSHPS61880.2024.10620845](https://doi.org/10.1109/INFOCOMWKSHPS61880.2024.10620845)]
- [28] Qi T, Wu FZ, Wu CH, Huang YF, Xie X. Uni-FedRec: A unified privacy-preserving news recommendation framework for model training and online serving. In: *Findings of the Association for Computational Linguistics: EMNLP 2021*. Punta: ACL, 2021. 1438–1448. [doi: [10.18653/v1/2021.findings-emnlp.124](https://doi.org/10.18653/v1/2021.findings-emnlp.124)]
- [29] Liang F, Pan WK, Ming Z. FedRec++: Lossless federated recommendation with explicit feedback. In: *Proc. of the 35th AAAI Conf. on Artificial Intelligence*. AAAI Press, 2021. 4224–4231. [doi: [10.1609/aaai.v35i5.16546](https://doi.org/10.1609/aaai.v35i5.16546)]
- [30] Liu SC, Ge YQ, Xu SY, Zhang YF, Marian A. Fairness-aware federated matrix factorization. In: *Proc. of the 16th ACM Conf. on Recommender Systems*. Seattle: ACM, 2022. 168–178. [doi: [10.1145/3523227.3546771](https://doi.org/10.1145/3523227.3546771)]
- [31] Lin YJ, Ren PJ, Chen ZM, Ren ZC, Yu DX, Ma J, de Rijke M, Cheng XZ. Meta matrix factorization for federated rating predictions. In: *Proc. of the 43rd Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval*. ACM, 2020: 981–990. [doi: [10.1145/3397271.3401081](https://doi.org/10.1145/3397271.3401081)]
- [32] Yang EY, Huang YF, Liang F, Pan WK, Ming Z. FCMF: Federated collective matrix factorization for heterogeneous collaborative filtering. *Knowledge-based Systems*, 2021, 220: 106946. [doi: [10.1016/j.knsys.2021.106946](https://doi.org/10.1016/j.knsys.2021.106946)]
- [33] Lin ZH, Pan WK, Yang Q, Ming Z. A generic federated recommendation framework via fake marks and secret sharing. *ACM Trans. on Information Systems*, 2022, 41(2): 40. [doi: [10.1145/3548456](https://doi.org/10.1145/3548456)]
- [34] He XN, Liao LZ, Zhang HW, Nie LQ, Hu X, Chua TS. Neural collaborative filtering. In: *Proc. of the 26th Int'l Conf. on World Wide Web*. Perth: Int'l World Wide Web Conferences Steering Committee, 2017. 173–182. DOI: [10.1145/3038912.3052569](https://doi.org/10.1145/3038912.3052569)
- [35] Zhang H, Wu B, Yuan XL, Pan SR, Tong HH, Pei J. Trustworthy graph neural networks: Aspects, methods, and trends. *Proc. of the IEEE*, 2024, 112(2): 97–139. [doi: [10.1109/JPROC.2024.3369017](https://doi.org/10.1109/JPROC.2024.3369017)]
- [36] Lei RZ, Wang PH, Zhao JZ, Lan L, Tao J, Deng C, Feng JL, Wang XD, Guan XH. Federated learning over coupled graphs. *IEEE Trans. on Parallel and Distributed Systems*, 2023, 34(4): 1159–1172. [doi: [10.1109/TPDS.2023.3240527](https://doi.org/10.1109/TPDS.2023.3240527)]
- [37] Wu CH, Wu FZ, Lyu L, Qi T, Huang YF, Xie X. A federated graph neural network framework for privacy-preserving personalization. *Nature Communications*, 2022, 13(1): 3091. [doi: [10.1038/s41467-022-30714-9](https://doi.org/10.1038/s41467-022-30714-9)]
- [38] Lin YJ, Han S, Mao HZ, Wang Y, Dally B. Deep gradient compression: Reducing the communication bandwidth for distributed training. In: *Proc. of the 2018 Int'l Conf. on Learning Representations*. OpenReview.net, 2018.
- [39] Rothchild D, Panda A, Ullah E, Ivkin N, Stoica I, Braverman V, Gonzalez J, Arora R. FetchSGD: Communication-efficient federated learning with sketching. In: *Proc. of the 37th Int'l Conf. on Machine Learning*. PMLR, 2020. 8253–8265.
- [40] Yun WJ, Kwak Y, Baek H, Jung S, Ji MY, Bennis M, Park J, Kim J. SlimFL: Federated learning with superposition coding over slimmable neural networks. *IEEE/ACM Trans. on Networking*, 2023, 31(6): 2499–2514. [doi: [10.1109/TNET.2022.3231864](https://doi.org/10.1109/TNET.2022.3231864)]
- [41] Chai D, Wang LY, Zhang JX, Yang L, Cai SW, Chen K, Yang Q. Practical lossless federated singular vector decomposition over billion-scale data. In: *Proc. of the 28th ACM SIGKDD Conf. on Knowledge Discovery and Data Mining*. Washington: ACM, 2022. 46–55. [doi: [10.1145/3534678.3539402](https://doi.org/10.1145/3534678.3539402)]
- [42] Chen HC, Vikalo H. Heterogeneity-guided client sampling: Towards fast and efficient Non-IID federated learning. In: *Proc. of the 38th Int'l Conf. on Neural Information Processing Systems*. Vancouver: Curran Associates Inc., 2024. 2093.
- [43] Nguyen J, Malik K, Zhan H, Yousefpour A, Rabbat M, Malek M, Huba D. Federated learning with buffered asynchronous aggregation. In: *Proc. of the 2022 Int'l Conf. on Artificial Intelligence and Statistics*. Valencia: PMLR, 2022. 3581–3607.
- [44] Tang ZH, Shi SH, Li B, Chu XW. GossipFL: A decentralized federated learning framework with sparsified and adaptive communication. *IEEE Trans. on Parallel and Distributed Systems*, 2023, 34(3): 909–922. [doi: [10.1109/TPDS.2022.3230938](https://doi.org/10.1109/TPDS.2022.3230938)]
- [45] Kim M, Saad W, Debbah M, Hong CS. SpaFL: Communication-efficient federated learning with sparse models and low computational overhead. In: *Proc. of the 38th Int'l Conf. on Neural Information Processing Systems*. ACM, 2024. 86500–86527.
- [46] Wang ZY, Xu HL, Xu Y, Jiang ZD, Liu JC, Chen S. FAST: Enhancing federated learning through adaptive data sampling and local training. *IEEE Trans. on Parallel and Distributed Systems*, 2024, 35(2): 221–236. [doi: [10.1109/TPDS.2023.3334398](https://doi.org/10.1109/TPDS.2023.3334398)]
- [47] Dwork C, McSherry F, Nissim K, Smith A. Calibrating noise to sensitivity in private data analysis. In: *Proc. of 3rd Theory of Cryptography Conf. on Theory of Cryptography*. New York: Springer, 2006. 265–284. [doi: [10.1007/11681878_14](https://doi.org/10.1007/11681878_14)]
- [48] Dwork C. Differential privacy. In: van Tilborg HCA, Jajodia S, eds. *Encyclopedia of Cryptography and Security*. New York: Springer, 2011. 338–340. [doi: [10.1007/978-1-4419-5906-5_752](https://doi.org/10.1007/978-1-4419-5906-5_752)]
- [49] Warner SL. Randomized response: A survey technique for eliminating evasive answer bias. *Journal of the American Statistical*

- Association, 1965, 60(309): 63–69. [doi: [10.1080/01621459.1965.10480775](https://doi.org/10.1080/01621459.1965.10480775)]
- [50] Li ZT, Wu XY, Pan WK, Ding YL, Wu ZH, Tan SQ, Xu Q, Yang Q, Ming Z. FedCORE: Federated learning for cross-organization recommendation ecosystem. *IEEE Trans. on Knowledge and Data Engineering*, 2024, 36(8): 3817–3831. [doi: [10.1109/TKDE.2024.3363505](https://doi.org/10.1109/TKDE.2024.3363505)]
- [51] Yang Q, Liu Y, Chen TJ, Tong YX. Federated machine learning: Concept and applications. *ACM Trans. on Intelligent Systems and Technology (TIST)*, 2019, 10(2): 12. [doi: [10.1145/3298981](https://doi.org/10.1145/3298981)]
- [52] Liang XX, Lin YQ, Fu HZ, Zhu L, Li XM. RSCFed: Random sampling consensus federated semi-supervised learning. In: *Proc. of the 2022 IEEE/CVF Conf. on Computer Vision and Pattern Recognition*. New Orleans: IEEE, 2022. 10154–10163. [doi: [10.1109/CVPR52688.2022.00991](https://doi.org/10.1109/CVPR52688.2022.00991)]
- [53] Li AR, Zhang L, Tan JT, Qin YX, Wang JH, Li XY. Sample-level data selection for federated learning. In: *Proc. of the 2021 IEEE Conf. on Computer Communications*. Vancouver: IEEE, 2021. 1–10. [doi: [10.1109/INFOCOM42981.2021.9488723](https://doi.org/10.1109/INFOCOM42981.2021.9488723)]
- [54] Wu Y, Zhang SC, Yu WC, Liu YC, Gu QQ, Zhou DW, Chen HF, Cheng W. Personalized federated learning under mixture of distributions. In: *Proc. of the 40th Int'l Conf. on Machine Learning*. Honolulu: JMLR, 2023. 1577.
- [55] Gao C, Huang C, Lin DS, Jin DP, Li Y. DPLCF: Differentially private local collaborative filtering. In: *Proc. of the 43rd Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval*. ACM, 2020. 961–970. [doi: [10.1145/3397271.3401053](https://doi.org/10.1145/3397271.3401053)]

作者简介

薛大暄, 博士生, 主要研究领域为隐私保护, 差分隐私, 深度学习.

杜宜霏, 硕士, 主要研究领域为联邦学习, 隐私保护, 深度学习.

陈红, 博士, 教授, 博士生导师, CCF 杰出会员, 主要研究领域为数据库技术, 新硬件平台下的高性能计算.

李翠平, 博士, 教授, 博士生导师, CCF 杰出会员, 主要研究领域为社交网络分析, 社会推荐, 大数据分析及挖掘.