

面向协作频谱感知的个性化差分隐私联邦学习方法

唐湘云¹, 康嘉文², 韩旭¹, 张焘³, 刘寅秋⁴, 孙庚⁵, 焦雨涛⁶

(1. 中央民族大学信息工程学院, 北京 100081; 2. 广东工业大学自动化学院, 广东 广州 510006;

3. 北京交通大学网络空间安全学院, 北京 100091; 4. 新加坡南洋理工大学, 新加坡 639798;

5. 吉林大学计算机科学与技术学院, 吉林 长春 130012; 6. 中国人民解放军陆军工程大学通信工程学院, 江苏 南京 210007)

摘 要: 针对协作频谱感知中数据非独立同分布 (Non-IID) 特性导致的模型性能下降问题, 提出了一种融合个性化差分隐私与重平衡分簇策略的联邦学习方案 (RebalFL)。该方案首先引入个性化差分隐私机制, 允许数据设置差异化隐私预算, 在保障隐私的同时减少噪声注入; 其次, 设计重平衡分簇策略, 构建数据分布均衡的客户端簇, 缓解模型漂移问题。实验结果表明, RebalFL 在 Non-IID 场景下显著优于现有差分隐私方法, 能有效提升频谱感知模型在隐私保护下的分类精度与鲁棒性。

关键词: 联邦学习; 个性化差分隐私; 频谱感知; 隐私保护

中图分类号: TN92

文献标志码: A

DOI: 10.11959/j.issn.1000-436x.2026067

Personalized differential privacy federated learning method for collaborative spectrum sensing

Tang Xiangyun¹, Kang Jiawen², Han Xu¹, Zhang Tao³, Liu Yinqiu⁴, Sun Geng⁵, Jiao Yutao⁶

1. School of Information Engineering, Minzu University of China, Beijing 100081, China

2. School of Automation, Guangdong University of Technology, Guangzhou 510006, China

3. School of Cyberspace Security, Beijing Jiaotong University, Beijing 100091, China

4. Nanyang Technological University, Singapore 639798, Singapore

5. College of Computer Science and Technology, Jilin University, Changchun 130012, China

6. College of Communications Engineering, PLA Army Engineering University, Nanjing 210007, China

Abstract: To address the degradation of model performance caused by data non-independent and identically distributed (Non-IID) characteristics in collaborative spectrum sensing, a federated learning scheme RebalFL was proposed, which integrated personalized differential privacy with a rebalancing clustering strategy. First, a personalized differential privacy mechanism that allowed heterogeneous privacy budgets for different data sources was introduced, thereby reducing noise injection while preserving privacy. Then, a rebalancing clustering strategy was designed to form client clusters with more balanced data distributions and mitigate model drift. Experimental results show that RebalFL outperforms existing differential privacy methods in Non-IID scenarios, substantially improving the classification accuracy and robustness of spectrum sensing models under privacy protection.

Keywords: federated learning, personalized differential privacy, spectrum sensing, privacy protection

收稿日期: 2025-12-01; 修回日期: 2026-03-04

通信作者: 韩旭, 25300517@muc.edu.cn

基金项目: 国家自然科学基金资助项目 (No.62572132, No.62572502, No.62571548); 国家密码基金资助项目 (No.2025NCSF02030); 国家自然科学基金青年科学基金资助项目 (C 类) (No.62302539, No.62402029); 北京市自然科学基金丰台联合基金资助项目 (No.L251041)

Foundation Items: The National Natural Science Foundation of China (No.62572132, No.62572502, No.62571548), The National Cryptography Foundation of China (No.2025NCSF02030), NSFC Youth Science Fund Project (Category C) (No.62302539, No.62402029), Beijing Natural Science Foundation Fengtai Joint Fund (No.L251041)

0 引言

无线电频谱是移动通信不可或缺的宝贵资源,随着通信技术的迅速发展和物联网设备的激增,全球通信产业对频谱资源的需求日益迫切。传统的频谱分配方式是静态地分配频谱,将无线频谱划分成若干固定的频段,由政府管理部门授权分配给用户独占使用。然而,基于授权的静态分配方式难以以为新增无线业务分配专用频段,并且实测调查发现,大部分授权频段的大多数时间处于空闲状态,频谱利用率很低。频谱资源不足和频谱利用率低的矛盾是目前亟待解决的问题^[1]。频谱共享被认为是提升频谱利用率和实现资源分配的重要途径,它允许不同用户在不同时间段共享同一段频谱数据。简单来说,它允许非授权用户(次用户(secondary user, SU))动态接入授权用户(主用户(primary user, PU))未使用的空闲频段,从而显著提升频谱利用率^[2]。

为实现动态频谱共享这一目标,必须构建一个具备感知、预测与决策能力的动态频谱接入(dynamic spectrum access, DSA)系统。如图1所示,典型的 DSA 系统由频谱感知、频谱预测与信道接入策略等功能模块构成。其中,频谱感知是系统运行的基础,其任务是检测主用户信号并识别可用空闲频段。由于后续的频谱预测与信道接入策略均依赖于频谱感知结果,故频谱感知的准确性直接决定了 DSA 系统的整体性能与频谱利用率。

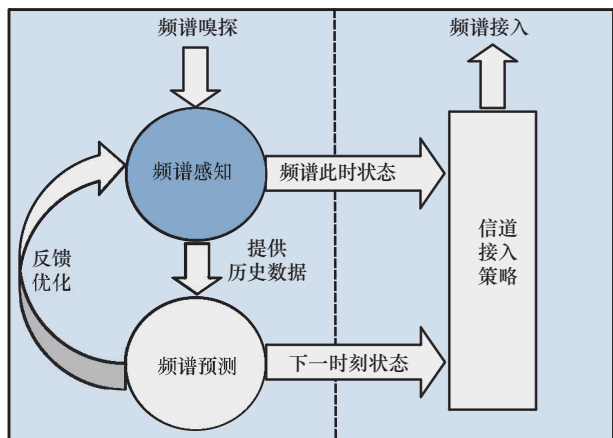


图1 动态频谱接入系统框架

频谱感知技术正在由单节点独立检测向多SU协作感知演进,旨在提升占用判决的可靠性。目前的协作频谱感知(cooperative spectrum sensing, CSS)主要有两种范式:一种是集中式协作频谱感知

(centralized CSS, CCSS),各SU把本地感知结果上传到融合中心,由融合中心统一决策;另一种是分布式协作频谱感知(distributed CSS, DCSS),各SU之间相互交换感知信息并形成一致判断^[3]。在上述两种范式中,协作的前提都是需要传输感知信息,而这些信息往往直接来自原始观测数据或由原始观测计算得到的能量统计量、局部判决或更细粒度的感知报告,因而在传输与共享过程中容易被反向推断出主用户的频谱占用状态等敏感信息,带来隐私泄露风险^[4]。

针对频谱感知数据传输带来的隐私安全挑战,联邦学习(federated learning, FL)作为一种新兴的分布式机器学习范式,为频谱感知模型训练提供了新思路。FL允许用户在本地设备上训练模型,仅将模型更新上传至中央服务器进行聚合,而不需要共享原始数据,在一定程度上保护了用户隐私。然而,研究表明,FL框架本身面临隐私泄露风险,如恶意攻击者可以通过成员推理攻击推断特定主用户是否参与训练^[5];通过模型逆向攻击重建主用户的频谱使用模式与地理位置信息^[6];通过梯度推断攻击从共享的模型梯度中精确反推出原始频谱样本数据^[7]。在频谱感知场景中,这些攻击可能导致主用户的通信频率、发射功率、活动规律等敏感信息被泄露,进而使其面临恶意干扰、精准定位追踪或非法占用频谱的风险。因此,为防止主用户隐私信息被恶意利用,有必要在联邦频谱感知中引入隐私保护机制。

在FL隐私保护技术中,差分隐私(differential privacy, DP)提供了一套严格的隐私保护体系,其核心优势在于通过数学证明确保隐私保护的强度,且隐私保护效果不依赖攻击者的背景知识,从而能够有效抵御多种背景知识下的推理攻击^[8]。然而,传统DP机制通常采用统一的隐私预算,对不同敏感程度的数据进行“一刀切”的噪声注入。在频谱感知场景中,这种处理方式不仅难以反映不同频谱状态在隐私风险上的差异,还可能直接影响频谱感知结果的可靠性。例如,当统一噪声水平设置偏大时,导致模型对频谱占用状态的判决不准确,容易将“繁忙”状态误判为“空闲”状态,从而导致次用户误入主用户正在使用的频段,增加干扰风险;当统一噪声水平设置偏小时,难以有效保护高敏感度样本所包含的信息,带来潜在的隐私泄露隐

患^[9]。因此,本文在联邦学习频谱感知框架中引入个性化差分隐私(personalized DP, PDP)机制,用户可根据频谱本地数据的敏感程度选择不同的隐私保护级别,一方面满足了用户层面对个性化隐私的需求,另一方面相较于统一DP预算的保护策略,可以减少全局的噪声注入以提高频谱感知模型的准确性^[10]。

然而,PDP会受到频谱感知数据非独立同分布(non-independent identically distributed, Non-IID)和标签依赖特性的制约。首先,不同区域的次用户所采集的本地数据存在显著差异,处于偏远地区或夜间时段的次用户更易获得“空闲”或“低占用”类样本,而位于城市中心或日间时段的用户则倾向于收集更多“繁忙”或“高占用”类样本,导致同一物理环境内的次用户数据近似独立同分布(independent identically distributed, IID),跨环境的次用户数据则呈现显著的Non-IID特性。在FL中,Non-IID数据首先会使各参与方的本地模型更新方向,导致全局模型收敛缓慢甚至发散。其次,频谱数据的隐私预算往往与标签紧密相关,以车联网为例,“高占用”信号蕴含更多车辆轨迹等敏感特征,易遭受位置推理攻击,通常需要分配较高的隐私预算进行强保护;“低占用”信号多为环境噪声,隐私风险低,分配较为宽松的隐私预算。这种基于标签的差异化隐私预算分配,实际上在Non-IID的基础上进一步引入了异质性的扰动噪声,将进一步诱发模型漂移现象,导致全局模型收敛更加困难。因此,亟须研究一种面向协作频谱感知的个性化差分隐私联邦学习方案,在保障差异化隐私需求的同时,有效缓解模型漂移并提升分类性能。

本文从优化频谱感知模型的方向出发,针对上述问题,提出了一种融合个性化差分隐私与重平衡分簇策略的联邦学习方案(RebalFL),旨在频谱数据Non-IID场景下,协同实现频谱安全共享与隐私保护。具体而言,本文首先引入PDP机制,允许不同客户端依据自身隐私需求与样本特性配置个性化隐私预算,在提供可靠隐私保护的同时尽可能减少噪声注入对模型性能的影响;其次设计一种重平衡分簇策略,通过将客户端数据划分为分布均衡的簇,并在簇内采用异步更新机制,以缓解数据Non-IID导致的模型漂移问题。本文的主要贡献如下。

1) 提出了一种融合个性化差分隐私与重平衡

分簇策略的联邦学习方案,应对频谱感知任务中数据标签不均衡与标签依赖特性所导致的模型漂移问题,实现了隐私保护与模型性能之间的动态平衡,显著提升了全局模型在Non-IID场景下的泛化能力。

2) 引入了PDP机制,系统性地防范训练过程中的隐私泄露风险,允许参与设备根据自身隐私需求灵活设置数据保护强度,在确保隐私安全的前提下,有效降低整体噪声的注入,从而提升模型实用性。

3) 设计了一种基于KL散度的重平衡分簇策略,将客户端划分为若干数据分布均衡的簇,使模型可在各平衡簇内执行分布式训练,从而有效缓解Non-IID数据带来的模型漂移问题,提升模型在异构环境下的鲁棒性。

4) 实验结果显示,RebalFL在多种Non-IID场景与差异化隐私预算设置下仍保持较高的模型准确率。在轻度Non-IID频谱任务中,RebalFL的F1值较PDPFL提升16.65%;在MNIST数据集的极端Non-IID场景中,RebalFL的模型准确率远高于其他3种基准方法,充分证明其能够有效缓解模型漂移问题并保持高准确率。

1 相关工作

本节将系统回顾与研究相关的现有工作,旨在梳理频谱感知从传统集中式与分布式协作频谱感知方法到基于FL的频谱感知再到隐私保护机制的研究脉络,奠定研究基础。

1.1 协作频谱感知技术现状

频谱感知的核心目标是从接收端观测信号中获取与主用户占用相关的可观测特征,进而判定目标频段的占用状态,为次用户的接入决策提供依据。随着无线环境的多径衰落与阴影效应加剧,单节点感知容易出现漏检或误警,因此研究重点已从单节点独立检测逐步转向协作频谱感知,并进一步发展为不需要固定融合中心、由各节点交换信息并迭代达成一致的分布式协作频谱感知,以提升感知结果的鲁棒性与可靠性。

因此形成了一系列基于统计理论的经典解析检测器,其共同特征是对人为设计的统计量特征进行检验。但是传统方法存在一些局限性,例如,基于能量阈值的能量检测(energy detection, ED)虽不

需要先验知识,但对噪声的干扰敏感,难以区分噪声与合法信号^[11];匹配滤波(matched filtering, MF)在已知先验序列并满足同步条件时最优,但对未知或异构主用户适用性有限^[12];循环平稳特征检测(cyclostationary feature detection, CFD)利用谱相关等循环特征抗噪,低信噪比(signal to noise ratio, SNR)有效但复杂度高且对采样与同步误差敏感^[13];协方差与特征值类方法利用相关结构缓解部分噪声的不确定性,但信息来源仍局限于有限阶统计相关^[14]。因此,传统方法能够刻画能量、相关与循环等显式可建模特征,但在多主用户叠加、强非平稳干扰、占用模式演化及硬件非理想等复杂场景下,少量统计量难以覆盖联合时频结构、时序依赖与非线性边界,导致鲁棒性与迁移能力受限^[15]。

为了在缺乏精确统计模型或信道环境剧烈变化的场景下提升感知精度,研究者近年来开始转向数据驱动的方案,现有研究广泛引入机器学习与深度学习方法。此类方法通过数据驱动的方式,从海量数据中自动提取高维判别性特征。例如,利用支持向量机(support vector machine, SVM)在高维空间构建最优超平面以实现二分类^[16];采用卷积神经网络(convolutional neural network, CNN)从时频图中提取局部空间纹理特征^[17];利用长短期记忆(long short-term memory, LSTM)网络捕捉信号序列的长时序依赖特征^[18]。总体来看,机器学习方法的优势主要体现在对复杂统计特性的刻画能力以及特征提取过程的自动化程度,这使其在统计模型难以精确刻画的场景下具备潜在优势。

然而,在基于协作频谱感知中,无论采用解析检测器还是机器学习、深度学习模型,协作过程都离不开感知信息的传输与共享,CCSS需要SU上传原始观测数据或判决统计量给融合中心,DCSS也需要邻居间交换细粒度感知报告。已有研究指出,这类感知数据或判决统计量可能被用于推断SU的位置等敏感信息,带来隐私泄露风险^[19]。

1.2 基于联邦学习的协作频谱感知研究现状

FL是一种分布式的机器学习架构,已初步应用于无线通信领域,为解决协作频谱感知任务带来的隐私问题,提出了更有前景的解决方案^[20]。Chen等^[21]针对单个频谱传感器在数据量和标注样本不足情况下训练性能受限的问题,提出了基于FL的多源协作频谱感知框架,通过跨节点知识共

享提升检测准确率。Song等^[22]提出了一种基于FL和多智能体深度强化学习算法的协作动态频谱接入技术,允许多用户在不共享训练数据的情况下协作实现目标系统优化。Shi等^[23]进一步验证了将传统的协作式频谱感知扩展至FL框架的可行性,为构建高效和隐私保护的分布式频谱监测系统提供了新思路。然而,在实际的无线通信系统中,用户设备和频谱数据往往是Non-IID的。为了解决这一问题,Li等^[24]在FL算法的优化过程中引入了一个近端项,并提出了一种近端FL算法。Wang等^[25]提出了一种基于归一化平均的FL算法来实现误差的快速收敛。尽管FL在一定程度上缓解了数据集中化带来的隐私风险,但其本身并未提供完善的隐私保护机制,在传输和聚合模型参数的过程中仍可能泄露敏感信息。

1.3 基于差分隐私的隐私保护联邦学习研究现状

DP作为一种严格的数学隐私保护框架,已广泛应用于FL,以缓解模型训练过程中的隐私泄露风险。传统DP机制通过在模型梯度或参数上传前注入噪声实现隐私保护^[26],从理论上可有效抵御多种攻击。然而,这类方法普遍采用统一的隐私预算,即对所有样本施加相同的隐私约束^[27-28]。这种“一刀切”的方式不仅忽视了不同用户对隐私保护的差异化需求,而且会对部分数据施加过强的隐私保护,从而引入过量噪声,降低全局FL模型性能。

为解决上述问题,研究者提出了PDP机制,允许参与者依据自身的隐私偏好设定个性化的隐私预算。根据个性化隐私预算分配的粒度,现有PDP机制可分为客户端级^[29]与记录级^[30]两类。前者如Liu等^[29]提出的客户端级PDP方案,为不同客户端分配不同隐私预算,以适应跨机构场景下多样化的隐私保护需求;后者则进一步细化至单条数据记录^[30],实现精确到每一条数据的个性化隐私控制,使不同数据具有不同的隐私预算。

在基于PDP的FL任务中,抽样机制发挥着关键作用。研究表明,引入随机抽样可有效防止隐私预算差异导致的数据灾难性遗忘问题^[31]。基于此,Boenisch等^[32]提出了一种基于二分搜索的采样概率求解算法,以在给定隐私预算下确定每条记录的最优抽样概率。然而,当隐私预算分布连续或范围较宽时,该算法的计算开销显著增加,难以适用于大规模FL场景。针对这一问题,Liu等^[30]进一步提出

了模拟曲线拟合 (simulation-curve fitting, SCF) 算法, 用于直接估计采样概率与隐私预算之间的非线性映射关系, 提高计算效率。

尽管 PDP 机制可以提升模型的整体效用, 但在 Non-IID 且隐私预算与标签相关的场景下存在显著局限性。现有研究通常假设客户端数据独立同分布, 没有考虑到客户端数据标签不均衡这一现实场景会导致模型漂移问题, 并且 PDP 机制可能进一步放大这一差异, 使模型分类偏向多数类, 从而削弱全局模型的整体泛化能力。因此, 本文提出了 RebalFL, 设计了重平衡分簇策略以有效缓解数据 Non-IID 和标签依赖特性引发的模型漂移问题。

2 预备知识

为便于后续的阅读与理解, 表1总结了本文使用的主要数学符号及其含义。

表1 主要数学符号及其含义

符号	含义
ε	差分隐私预算
δ	差分隐私松弛参数
ω	本地模型参数
T	总通信轮数
τ	本地迭代次数
q	采样概率
σ	高斯噪声标准差
$\text{KL}(\cdot \ \cdot)$	KL 散度
S_0	中央服务器
$S = \{S_1, S_2, \dots, S_K\}$	边缘服务器集合
$\mathcal{C} = \{C_1, C_2, \dots, C_I\}$	客户端集合

2.1 基于机器学习的频谱感知模型

频谱感知的核心任务是根据次用户接收端的观测信号判断目标频段是否存在主用户信号, 通常可将该问题建模为二元假设检验。设在一个观测窗内的离散接收序列为 $\{r[k]\}_{k=1}^N$, 则二元假设检验可表示为

$$\begin{cases} \mathcal{H}_0: r[k] = n[k] \\ \mathcal{H}_1: r[k] = s[k] + n[k] \end{cases} \quad (1)$$

其中, $n[k]$ 表示加性噪声, $s[k]$ 表示主用户信号。频谱感知通过构造统计量 $T[r]$ 并与阈值 λ 比较完成判决, 当 $T[r] \geq \lambda$ 时, 判为“忙碌” $\mathcal{H}_1$; 否则

判为“空闲” $\mathcal{H}_0$ 。对应的检测性能通常用检测概率与虚警概率表示为

$$P_d = \Pr(T(r) \geq \lambda | \mathcal{H}_1), \quad P_f = \Pr(T(r) \geq \lambda | \mathcal{H}_0) \quad (2)$$

在传统解析检测器中, 常见统计量包括能量统计量、循环平稳特征统计量等。以能量统计量为例, 可由接收信号在频域或时域的能量聚合得到, 如式(3)所示。

$$T_E(r) = \frac{1}{N} \sum_{k=1}^N |r[k]|^2 \quad (3)$$

该类方法的关键是统计量的设计通常依赖信号与噪声统计特性的先验假设, 且阈值设置对噪声不确定性较敏感。

为减少对精确统计模型与人工特征设计的依赖, 近年来频谱感知广泛引入机器学习方法, 将“是否占用”判决转化为监督分类问题。具体而言, 首先从观测序列 r 中构造特征向量 x 的特征可取能量、功率谱密度、时频图等。随后用带标签的数据集 $\{(x_i, y_i)\}_{i=1}^M$ 训练分类器 $f_\theta(\cdot)$, 其中 $y_i = \{0, 1, \dots, i\}$ 表示数据标签。训练过程可抽象为最小化经验风险

$$\min_{\theta} \frac{1}{M} \sum_{i=1}^M \ell(f_\theta(x_i), y_i) \quad (4)$$

其中, $\ell(\cdot)$ 表示分类损失函数。在推理阶段, 对新观测值 r^* 计算 x^* 并输出 $\hat{y} = f_\theta(x^*)$ 作为占用判决。

2.2 基于联邦学习的协作频谱感知框架

频谱感知是 DSA 系统的核心机制, 用于探测某特定频谱是否被主用户占用, 从而为次用户提供可靠的频谱接入机会, 提高频谱资源利用率^[33]。然而, 高精度的频谱感知模型往往依赖大量的历史数据, 而单个节点往往难以独立收集。为克服这一数据孤岛问题, FL 作为一种分布式机器学习范式, 为频谱感知提供了新的训练模式。

基于 FL 的频谱感知建模包含一个中央服务器和 I 个次用户客户端 $\mathcal{C} = \{C_1, C_2, \dots, C_I\}$ 。在频谱探测阶段, 各个次用户客户端收集得到频谱感知数据集 $D = \{D_1, D_2, \dots, D_I\}$, 并保存在本地。在训练阶段, 中央服务器负责协调所有次用户客户端共同训练一个全局的频谱感知模型 $\omega^* \in \mathbb{R}^d$ 。整个训练过程遵循 FL 流程, 包含以下 4 个核心步骤。

1) 模型下发: 在第 t 轮训练开始时, 中央服务器将当前的全局模型参数 ω' 广播至所有参与的次用户客户端。

2) 本地训练: 每个次用户客户端 C_i 接收模型参数后, 利用其本地数据集 D_i 独立进行模型训练, 通过梯度下降更新本地模型参数 ω_i^t 。

$$\omega_i^{t+1} = \omega_i^t - \eta l_i(\omega_i^t; D_i) \quad (5)$$

其中, η 为学习率, ω_i^t 为次用户客户端 C_i 的本地模型, $l_i(\omega_i^t; D_i)$ 为次用户客户端 C_i 上的局部损失函数。

3) 模型上传: 各个次用户客户端在完成本地训练后, 将更新的模型参数 ω_i^{t+1} 上传至中央服务器。

4) 全局聚合: 中央服务器在接收到来自各个次用户客户端的模型更新后, 采用加权平均的方式进行全局模型聚合, 得到新一轮的全局模型。

$$\omega^{t+1} = \sum_{i=1}^I \frac{n_i}{N} \omega_i^{t+1} \quad (6)$$

其中, n_i 为次用户客户端 C_i 的本地数据量, N 为所有参与客户端的总样本数。

循环执行上述过程, 直至模型收敛或达到预设训练轮数, 输出最终的全局频谱感知模型 ω^* , 由中央服务器广播到各次用户客户端投入使用。

2.3 个性化差分隐私机制

DP 是一种基于数据驱动的隐私保护方法, 结合 DP 可以进一步降低 FL 的隐私泄露风险。PDP 是传统 DP 的一种变体, 旨在针对数据集中的每条数据进行不同程度的隐私保护, 防御攻击者从共享的模型更新中推断敏感信息。为此, 本节首先回顾 DP 的基本概念, 进而介绍 PDP 以及本文使用的噪声机制。

定义 1 (ϵ, δ) -DP。随机机制 A 与所有可能输出结果构成的集合 $o \in \text{range}(A)$, 在相邻数据集 D 和 D' 上, 对任意的输出查询, 若满足式(7)所示条件, 则称随机机制 A 满足 (ϵ, δ) -DP^[34]。

$$\Pr[A(D) \in o] \leq e^\epsilon \Pr[A(D') \in o] + \delta \quad (7)$$

其中, $\Pr[A(D) \in o]$ 为随机机制 A 在数据集 D 上输出落在集合 o 内的概率; ϵ 为隐私预算, 作为隐私保护水平的度量指标, 隐私预算越大, 表示保护水平越低, 隐私预算越小, 表示保护水平越高; δ 为松弛参数, 表示随机机制 A 不满足 (ϵ, δ) -DP 的概率。

定义 2 (ϵ_{ij}, δ) -PDP。随机机制 A 与其所有可能输出结果构成的集合 $o \in \text{range}(A)$ 在相邻数据集 D 和 D_{-i} 上, 对任意的输出查询, 若满足式(8)所示条件, 则称随机机制 A 满足 (ϵ_{ij}, δ) -PDP^[35]。

$$\Pr[A(D) \in o] \leq e^{\epsilon_{ij}} \Pr[A(D_{-i}) \in o] + \delta \quad (8)$$

定义 3 高斯机制。通过向查询结果中添加服从正态分布的噪声来保护隐私。对于给定查询函数 f , 在数据集 D 上的输出可表示为

$$M(D) = f(D) + N(0, \sigma^2) \quad (9)$$

其中, $N(0, \sigma^2)$ 为噪声项, 具有均值为 0、方差为 σ^2 的高斯分布。为了满足 (ϵ, δ) -PDP, 需要选取合适的 σ , 常见的一种设置方式为

$$\sigma = \sqrt{2 \ln\left(\frac{2}{\delta}\right) \frac{\Delta f}{\epsilon}} \quad (10)$$

其中, Δf 为函数的敏感度。

3 问题定义

3.1 系统模型

本文考虑一个多区域无线频谱感知场景。每个区域内存在大量 PU, 根据业务不同, 被授权不同频段进行通信, 持续产生真实的频谱信号。为实现对频谱占用状态的感知, 部署若干频谱嗅探节点作为次用户客户端, 负责在本地进行频谱感知模型训练。所有次用户通过无线网络与云端的中央服务器通信, 中央服务器负责协调 FL 的全局感知模型训练与分发。

基于 FL 的协作频谱感知系统架构如图 2 所示。该系统架构包含 3 类实体: 主用户、次用户客户端和中央服务器。各实体的具体功能归纳如下。

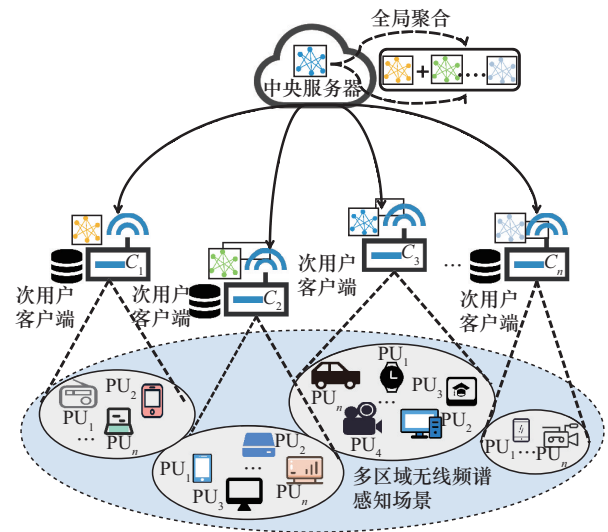


图2 基于联邦学习的协作频谱感知系统架构

1) 主用户。主用户是频段授权的持有者, 其活动数据被次用户客户端通过频谱感知技术采集与

标注,作为频谱感知建模的原始数据集。

2) 次用户客户端。每个次用户客户端利用本地传感器采集频谱感知数据,在FL过程中,接收来自中央服务器的全局模型,在本地执行训练任务,仅将本地模型更新上传至中央服务器,确保原始数据不出本地。

3) 中央服务器。中央服务器负责初始化与聚合全局模型。在每轮训练迭代中,中央服务器聚合来自多个次用户客户端的本地模型更新,通过联邦平均等算法更新全局模型。

3.2 威胁模型

本文假设次用户客户端与中央服务器均为诚实但好奇的参与者,这意味着它们会严格遵循预设的训练协议,但可能试图通过共享信息进行隐私推理^[36]。次用户客户端在FL中能得到每一轮的全局模型更新,可能通过分析其获得的全局模型,尝试推断其他客户端的敏感数据,发起成员推断攻击^[37]。中央服务器作为全局模型的维护者,能够接触所有客户端的模型更新,可能通过分析这些模型更新尝试重构或推断特定客户端的本地数据^[38]。

3.3 设计目标

本文面向协作式的频谱感知场景,设计目标具体如下。

1) 隐私保护。针对不同主用户对隐私保护强度的差异化需求,允许不同数据设置满足偏好的隐私预算。方案需从理论上证明能够满足PDP,提供严格隐私保障。

2) 缓解Non-IID场景下的模型漂移问题。方案能有效防止因数据分布不均衡或隐私保护强度差异而导致的模型分类偏向特定类别的模型漂移问题。

3) 隐私保护下模型准确率无损失。在引入PDP机制后需确保全局模型不会因噪声干扰而出现全局模型的精确度大幅下降,其最终性能应优于传统DP方法,确保频谱感知模型输出的可靠性。

4 方案设计

为实现上述目标,本文提出了一种融合个性化差分隐私与重平衡分簇策略的联邦学习方案(RebalFL)。RebalFL在提供PDP机制的前提下,能有效缓解因数据Non-IID特性导致的模型漂移问题,最终通过FL训练后的全局模型将拥有较高的模型准确率。

4.1 方案概述

RebalFL的整体方案概述如图3所示,包含3个核心实体:中央服务器、边缘服务器和次用户客户端。

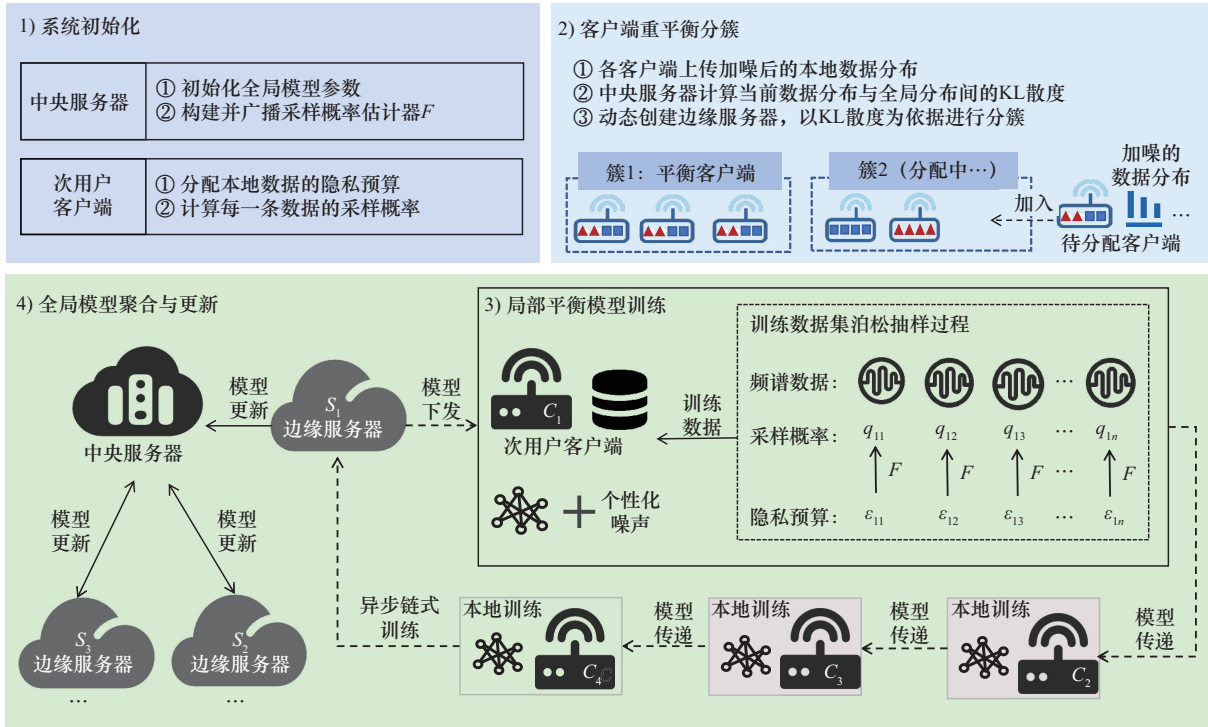


图3 RebalFL的整体方案概述

RebalFL 方案的执行流程是一个迭代的过程,可分为以下 4 个阶段。

1) 系统初始化。中央服务器初始化全局模型参数,构建并广播采样概率估计器,各个次用户客户端完成本地数据的隐私预算分配与采样概率计算。

2) 客户端重平衡分簇。中央服务器基于客户端的数据分布并通过 KL 散度计算,构建多个数据分布平衡的客户端簇,每个簇引入一个边缘服务器进行管理。

3) 局部平衡模型训练。各边缘服务器在其簇内,组织客户端以链式顺序进行模型的异步链式训练与传递。客户端在本地训练中应用 PDP 保护机制,最终由边缘服务器提交局部平衡模型至中央服务器。

4) 全局模型聚合与更新。中央服务器收集所有边缘服务器提交的模型更新,通过联邦平均算法聚合生成新一代的全局模型。随后,返回第一阶段,循环直至模型收敛或达到预设终止条件。

4.2 系统初始化

系统初始化阶段由中央服务器与次用户客户端协同完成,为后续实现隐私保护与高效训练奠定基础。该阶段主要包括中央服务器与客户端初始化两部分。

在中央服务器初始化阶段,首先,对全局频谱感知模型的参数进行初始化;随后,中央服务器通过模拟曲线拟合算法构建采样概率估计器,该过程借助数值模拟方法,探究在不同采样概率下,完成完整联邦训练后单个数据记录所能达到的实际隐私成本上界。模拟结果显示,隐私预算与采样概率之间呈指数函数关系 $\varepsilon \approx e^{aq+b} + c$ 。基于此关系,通过模拟曲线拟合算法得到采样概率估计器 $F(\varepsilon) = q = \frac{\ln(\varepsilon - c) - b}{a}$, 其中 $a \neq 0$, $\varepsilon \in (0, 10]$, $c \in (0, \varepsilon)$ 。该采样概率估计器能够依据任意给定的隐私预算,直接计算出相应的最优采样概率。

在客户端初始化阶段,首先,各客户端 C_i 根据数据敏感度或具体隐私需求,为其本地数据集中的每一条频谱数据 d_{ij} 分配不同的隐私预算 ε_{ij} ; 随后,客户端使用采样概率估计器为其本地数据集中的每一条频谱数据计算对应的最优采样概率 q_{ij} 。

4.3 客户端重平衡分簇

由于不同地理位置的次用户客户端所处信号环

境存在显著差异,其感知数据标签也存在差异。为缓解由此引起的数据分布不均衡问题,本节通过客户端重平衡分簇策略,对每一个趋于全局平衡分布的簇引入一个边缘服务器进行管理,分为以下两个阶段。

在客户端数据上传阶段,各客户端首先在本地上统计其数据集的类别标签和数量,随后对这些统计信息进行 DP 加噪,最后将加噪后的数据分布上传至中央服务器。由此,中央服务器获得的是经噪声扰动的数据分布信息,而无法获取任何客户端本地数据的具体内容,从而能够在实现数据分布可见的同时确保原始频谱感知数据的隐私不被泄露。

在客户端重平衡分簇阶段,通过边缘服务器调度实现客户端的重平衡分簇,缓解各客户端数据分布不均衡带来的模型漂移问题。为此,系统首先引入若干边缘服务器,并采用贪婪策略依次为其分配客户端。具体而言,每个边缘服务器在遍历所有未分配客户端数据分布后,选择使数据分布最接近均匀分布的客户端,该最优选择通过最小化当前数据与均匀分布之间的 KL 散度来实现(如算法 1 第 4 行所示)。当某个边缘服务器达到其客户端容量上限时,中央服务器将创建新的边缘服务器并继续执行该过程,直至所有客户端均完成分配。

算法 1 基于 KL 散度的客户端重平衡分簇

输入 客户端集合 $\mathcal{C} = \{C_1, C_2, \dots, C_I\}$, 边缘服务器集合 $S = \emptyset$, 全局平衡分布 P_u

输出 分配好客户端的边缘服务器集合 $S = \{S_1, S_2, \dots, S_K\}$

1) repeat

2) 创建边缘服务器 S_k

3) for $|\mathcal{C}| > 0$ and $|S_k| < \gamma$ do

4) $S_k \leftarrow \operatorname{argmin}_i D_{\text{KL}}(P_{S_k} + P_{C_i} \| P_u)$, $C_i \in \mathcal{C}$

5) 边缘服务器 S_k 加入客户端 C_i

6) 从 $\mathcal{C}$ 中移除客户端 C_i

7) end for

8) $S \leftarrow S \cup S_k$

9) return S

接下来对算法 1 进行时间复杂度分析。设客户端总数为 n , 数据标签类别数为 L 。在重平衡分簇过程中,寻找最优匹配需遍历候选集并计算 KL 散

度, 计算总次数为 $\sum_{i=1}^n (n-i+1) = \frac{n(n+1)}{2}$, 单

次KL散度运算的时间复杂度为 $O(L)$, 因此算法1的总时间复杂度为 $O(n^2L)$ 。由于 L 在实际频谱感知场景中通常为极小常数, 且该算法仅在初始化阶段执行一次, 因此整体方案在面对大规模客户端接入时仍展现出良好的可扩展性。

当前算法采用局部最优的贪婪选择策略, 在极端数据分布下可能难以获得全局最优的分簇结果, 具体表现为优先构建的平衡簇可能提前消耗关键类别样本, 使后续客户端难以匹配理想的均衡分布。未来工作可考虑引入启发式搜索或群体智能优化方法, 以提升分簇结果的全局最优性与鲁棒性。

4.4 局部平衡模型训练

在完成分簇后, 各边缘服务器在其管理的簇内组织客户端进行模型训练。模型从首个客户端开始, 按顺序在客户端间传递与更新, 最后由边缘服务器收集所有更新并进行局部聚合。其中, 重点介绍基于PDP的客户端本地训练流程, 主要包括以下4个步骤。

步骤1 模型下发。边缘服务器将当前簇的局部模型 ω_i^r 下发至第一个训练的客户端。

步骤2 记录级泊松采样。在开始训练前, 客户端 C_i 对本地每一条数据 d_{ij} 按记录采样概率 q_{ij} 独立执行泊松抽样 $\beta_{ij} \sim \text{Ber}(q_{ij})$, 得到本轮参与训练的小批量训练集 $D_i = \{d_{ij} | \beta_{ij} = 1\}$ 。通过泊松抽样机制, 可以利用隐私放大效应降低隐私开销。

步骤3 本地模型训练与更新。客户端在模型训练过程中对梯度添加高斯噪声 $N(0, \sigma^2)$, 并持续跟踪每条数据的累计隐私消耗 ρ_{ij} 。当某条记录的隐私预算耗尽时, 该记录将不再参与后续训练。训练完成后, 客户端将更新后的本地模型 ω_i^{r+1} 传递给下一个客户端或边缘服务器。

步骤4 模型提交。当边缘服务器拿到最后一个客户端提交的模型更新后, 若没有达到局部迭代次数上限 τ , 则重新交给第一个客户端执行步骤3; 反之, 提交到中央服务器, 完成一轮训练, 返回步骤1, 直到模型收敛或达到全局训练轮数 T 。

在训练完成后, RebalFL能够确保在整个FL过程中, 每条数据的累计隐私消耗 ρ_{ij} 不超过其预设

隐私预算 ε_{ij} , 从而为系统中的每一个数据点提供可证明的 $(\varepsilon_{ij}, \delta)$ -PDP 保护。

5 安全性分析

本文旨在确保 RebalFL 在协作频谱感知过程中满足 $(\varepsilon_{ij}, \delta)$ -PDP, 防御“诚实但好奇”的中央服务器及次用户客户端发起的推理攻击。具体而言, 即便“诚实但好奇”的中央服务器及次用户客户端能够获取全局模型参数或梯度更新信息, 也无法推断任意本地数据的存在性。

首先进行敏感度分析, 保证 RebalFL 在任意一对相邻数据集 D'_i 与 D_i 上的输出分布在统计上是不可区分的, 从而掩盖单条数据记录的存在性。为了限制单条记录对模型更新的影响, RebalFL 在本地训练阶段引入了梯度裁剪机制。

假设次用户客户端 C_i 拥有相邻数据集 D'_i 与 D_i , $f(D_i)$ 为客户端在数据集 D_i 上计算得到的梯度更新函数。在执行梯度下降前, 将每条样本的梯度 g 的 L_2 范数限制在阈值 C 以内, 即 $\|g\|_2 \leq C$ 。根据梯度裁剪机制的性质, 随机梯度下降函数的 L_2 敏感度 Δs 上界推导为

$$\Delta s = \max_{D_i, D'_i} \|f(D_i) - f(D'_i)\|_2 = \eta \max \|g_i - g'_i\| = \frac{\eta}{|D_i|} 2C \leq C \quad (11)$$

以上结果确保了任意单条数据的加入或移除对模型参数更新量的影响被严格限制在阈值 C 内, 这表明 RebalFL 的训练过程对单条数据点的变化具有稳定性, 即模型更新对任意记录的依赖被严格限制, 从而满足 DP 的敏感度要求。

其次进行隐私性分析, 基于 Rényi 差分隐私 (Rényi differential privacy, RDP) 框架对 RebalFL 的隐私损失上界进行推导。通过分析泊松采样带来的隐私放大效应, 计算多轮训练的 PDP 隐私累计消耗, 以及 RDP 与 PDP 的转换定理, 最终推导出 RebalFL 满足 $(\varepsilon_{ij}, \delta)$ -PDP。

RebalFL 采用泊松采样高斯机制实现 DP。对于记录 d_{ij} , 设其采样概率为 q_{ij} , 噪声标准差为 σ 。RDP 利用 Rényi 散度量化相邻数据集在随机算法输出分布上的差异, 能够对隐私泄露进行更精细和严谨的控制。对于 $\alpha > 1$ 阶的 Rényi 散度, 单次本地迭代对记录 d_{ij} 造成的隐私损失 $\rho_{\text{iter}}(a, q_{ij})$ 满足式(12)

所示上界^[39]。

$$\rho_{\text{iter}}(\alpha, q_{ij}) \leq \frac{1}{\alpha - 1} \text{lb} \{ (1 - q_{ij})^{\alpha-1} (\alpha q_{ij} - q_{ij} + 1) + \left(\frac{\alpha}{\ell} \right) (1 - q_{ij})^{\alpha-\ell} q_{ij}^{\ell} e^{(\ell-1)\rho(\ell)} \} \quad (12)$$

其中, $\rho(a) = \frac{a}{2\sigma^2}$ 。由式(12)可以看出, 采样概率 q_{ij} 越低, 单步隐私损失 ρ_{iter} 越小, 体现了采样带来的隐私放大效应。

在 FL 中, 模型训练涉及多轮通信与本地迭代。假设总通信轮数为 T , 每轮包含 τ 次本地迭代。根据 RDP 的自适应序列组合定理^[40], 总隐私损失等于各步骤隐私损失的线性组合。因此, 对于记录 d_{ij} , 在完成所有训练轮次后, 其累计的 RDP 隐私损失 $\rho_{\text{total}}(a)$ 为

$$\rho_{\text{total}}(a) = \sum_{t=1}^T \sum_{r=1}^{\tau} \rho_{\text{iter}}^{(t,r)}(a) = T\tau \rho_{\text{iter}}(a, q_{ij}) \quad (13)$$

最后, 利用 RDP 与 DP 的转换引理, 对于给定的松弛项 $\delta > 0$, 可以将累计的 RDP 参数转换为 (ϵ_{ij}, δ) -PDP, 表示为

$$\epsilon_{ij} = \min_{a > 1} \left[\rho_{\text{total}}(a) + \frac{\text{lb}\left(\frac{1}{\delta}\right)}{a-1} \right] \quad (14)$$

通过上述推导可证明, RebalFL 可以精确控制记录 d_{ij} 在 $T\tau$ 轮训练中的累计隐私损失 ρ_{total} , 确保最终的隐私消耗严格受限于用户设定的 PDP 预算 ϵ_{ij} 。因此, RebalFL 满足 (ϵ_{ij}, δ) -PDP。

6 实验评估

本节介绍 RebalFL 的实验设计及结果分析, 重点评估 RebalFL 在进行个性化隐私保护的同时, 能否有效缓解由 Non-IID 数据特性引发的模型漂移问题, 从而显著提升全局模型在联邦训练过程中的收敛稳定性与最终预测准确率。

6.1 实验设置

6.1.1 数据集与模型介绍

为全面评估 RebalFL 的有效性, 本文基于 MNIST 与模拟频谱感知数据集开展实验。具体而言, 首先, 基于 MNIST 数据集模拟现实频谱占用度三分类任务。其次, 基于模拟生成的频谱感知数据集, 进行了更贴近实际频谱感知场景的频谱可用性二分类任务。实验用到的数据集与模型介绍

如下。

1) MNIST 数据集。选取 MNIST 数据集中标签为 0、1、2 的样本, 分别映射频谱感知中的“低占用”“中度占用”与“高度占用”样本。每个样本为 28×28 的灰度图像。该子集共包含 18 623 个训练样本与 3 147 个测试样本, 类别分布基本均衡, 适用于验证算法在通用分类任务中的基础性能。

2) 模拟频谱感知数据集。为确保模拟实验在认知无线电场景下的有效性, 数据集生成方法参考文献^[41]。该数据集的构建逻辑从二元假设检验出发, 通过提取接收信号的能量统计量作为核心分类特征, 旨在模拟实际系统中能量检测器对信道状态(空闲或繁忙)的识别判决过程, 其中标签类别 0 表示空闲状态, 标签类别 1 表示繁忙状态。训练集包含每类各 5 000 个样本, 测试集包含 1 500 个空闲样本与 900 个繁忙样本, 训练集与测试集总体划分比例约为 8:2。

3) 模型介绍。实验使用 CNN 作为核心分类模型, 其中卷积层和池化层用于对输入特征进行提取。在每次全局通信迭代时, 设置全局通信迭代轮数为 15 轮, 客户端本地训练的迭代次数为 50, 用于本地更新的 batch_size 为 128, 学习率为 0.1。

6.1.2 隐私预算分配策略

为评估 RebalFL 在不同隐私预算分配策略下的性能表现, 针对不同任务设计了多种隐私预算分配场景。在 MNIST 三分类任务中, 为标签 0、1 和 2 样本分别分配隐私预算 ϵ_1 、 ϵ_2 和 ϵ_3 , 并设计了 3 种分配方式。

1) 固定预算。每个类别分配特定且不变的隐私预算值, 即 $\epsilon_1 = 0.1$ 、 $\epsilon_2 = 1.0$ 和 $\epsilon_3 = 5.0$ 。

2) 多分布预算。3 个类别的隐私预算分别从 3 个不同的正态分布中采样, 即 $N_1(0.1, 0.01)$ 、 $N_2(1.0, 0.05)$ 和 $N_3(5.0, 0.5)$, 以模拟各客户端隐私需求存在一定波动的情形。

3) 长尾分布预算。所有样本的隐私预算整体服从一个帕累托分布, 其形状参数值为 1.0, 下限为 0.1, 模拟大多数样本隐私需求高、少数样本隐私需求低的长尾分布特性。

在模拟频谱二分类任务中, 采用固定预算分配方式: 为类别 0 赋予宽松的隐私预算 $\epsilon_1 = 0.5$, 为类别 1 设置严格的隐私预算 $\epsilon_2 = 0.05$, 以此评估该方法在隐私保护强度差异显著场景下的模型性能。

6.1.3 基准方法

为全面评估所提方法的性能,本文选取了3种具有代表性的基线方法进行对比分析。 $\text{FedAvg}^{[42]}$ 作为经典的联邦平均算法,不引入任何隐私保护机制,为实验提供了性能参考。 $\text{PDPFL}^{[30]}$ 使用PDP机制,允许不同数据根据自身需求设置差异化隐私预算。 $\text{MidFL}^{[43]}$ 采用统一的DP-FedAvg机制,其隐私预算取所有预算的平均值,对所有客户端的数据施加同等强度的隐私保护。

6.2 隐私保护下的RebalFL模型准确率评估

本节介绍RebalFL与3种基准方法在MNIST三分类任务中的实验设置及结果分析。实验将在模型准确率和运行时间两方面进行系统评估,重点考察RebalFL在不同次用户客户端数据异构场景及隐私预算分配方式下是否能保持较高的模型准确率,所有取值均为3次实验的平均值。

设置30个次用户客户端,并构建3种数据异构场景:在IID场景下,所有次用户客户端均匀持有3类标签的样本,构成理想的数据分布;在轻度异构(Non-IID1)场景下,10个次用户客户端仅包含单一标签样本,10个次用户客户端包含两种标签样本,其余10个次用户客户端包含全部3类标签样本,以模拟现实中常见的分布不均情况;在重度异构(Non-IID2)场景下,每个次用户客户端仅包含其中一类标签样本。

6.2.1 IID场景下的性能对比

本节在IID场景下,分别在固定预算、多分布预算和长尾分布预算3类隐私预算设置下对各方法进行了评估。实验结果如表2与图4所示,其中,本文方案的结果以加粗字体突出显示。RebalFL在所有设置下均保持了较稳定的训练曲线和较高的模型准确率,尽管在固定预算与多分布预算设置下,其最终模型准确率较低,但RebalFL的整体性能仍然较高,且在极端的长尾分布预算设置下表现优异。

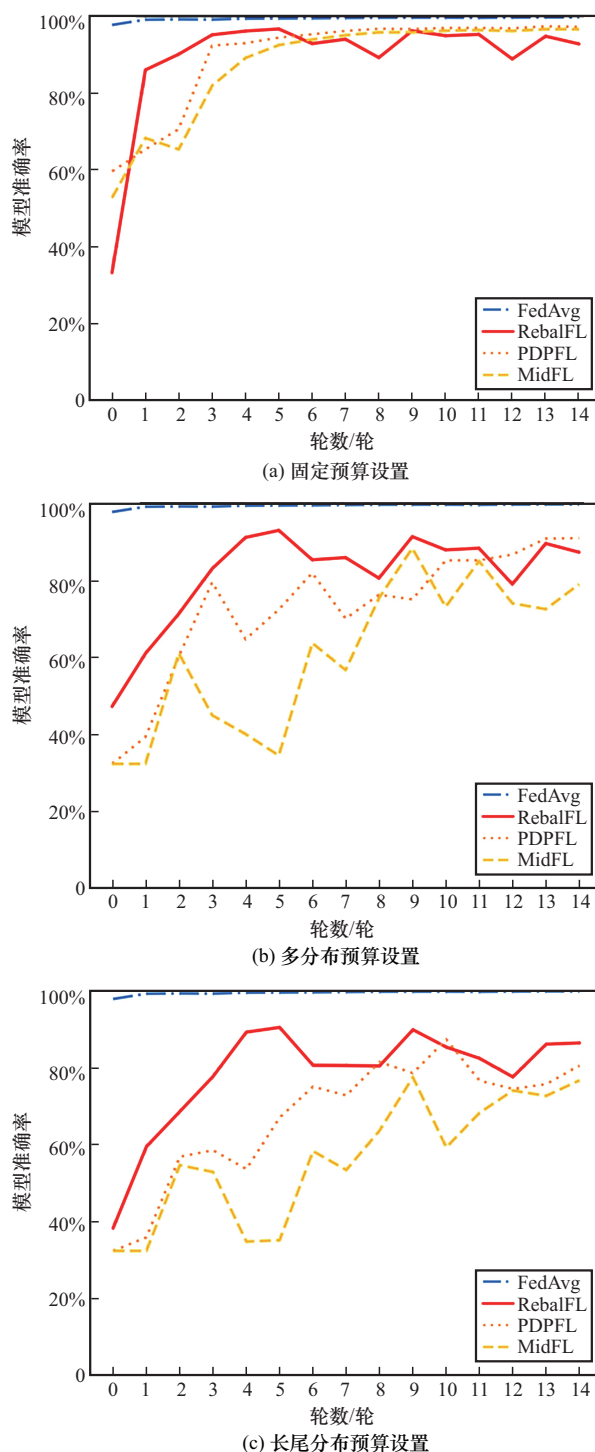


图4 IID场景下各方法的模型准确率对比

表2

IID场景下各方法的性能对比

隐私预算设置	$\text{FedAvg}^{[42]}$		$\text{MidFL}^{[43]}$		$\text{PDPFL}^{[30]}$		RebalFL (本文方案)	
	模型准确率	运行时间/s	模型准确率	运行时间/s	模型准确率	运行时间/s	模型准确率	运行时间/s
固定预算	99.65%	218	96.44%	482	97.15%	630	92.13%	620
多分布预算	99.65%	218	75.18%	243	89.47%	276	85.28%	275
长尾分布预算	99.65%	218	74.40%	230	76.84%	259	83.27%	256

在极端的长尾分布预算设置下, FedAvg 由于未引入任何噪声, 因此模型准确率最高, 其次是 RebalFL 取得了 83.27% 的模型准确率, 显著高于 PDPFL 的 76.84% 与 MidFL 的 74.40%。这一结果主要源于 RebalFL 的个性化差分隐私机制与重平衡分簇策略。首先, 个性化差分隐私机制相对于统一差分隐私机制能减少全局噪声注入量, 因此 RebalFL 与 PDPFL 的模型准确率高出 MidFL。其次, 重平衡分簇策略能有效缓解由预算差异导致的低预算样本参与度不均, 使高预算与低预算数据对模型的贡献更为均衡, 因此 RebalFL 的模型准确率显著高于 PDPFL。

从运行时间来看, 在固定预算和多分布预算设置下, RebalFL 的训练时间均高于 FedAvg 和 MidFL, 这主要源于 PDP 机制与重平衡分簇策略所带来的额外计算开销。然而, 相较于 RebalFL 在模型精度和稳定性方面的显著提升, 这部分额外计算开销处于可接受范围。

6.2.2 轻度异构 Non-IID1 场景下的性能对比

本节在轻度异构 Non-IID1 场景下, 分别在固定预算、多分布预算和长尾分布预算 3 类隐私设置中对各方法进行了实验评估。实验结果如表 3 与图 5 所示, 其中, 本文方案的结果以加粗字体突出显示。所有隐私保护方法的模型准确率相较于 IID 场景均出现不同程度下降, 说明数据轻度异构已经加剧了模型漂移问题, 进而导致模型准确率下降。尽管 FedAvg 由于没有引入噪声而仍保持最高模型准确率, 但其训练前期出现明显振荡。相比之下, RebalFL 在所有隐私预算设置下均呈现更平稳的提升曲线, 并在所有隐私方法中保持最优的整体性能。

在固定预算设置下, RebalFL 达到 95.36% 模型准确率, 优于 MidFL 的 89.25% 与 PDPFL 的 93.02%。在多分布预算设置下, MidFL 和 PDPFL 分别下降至 69.48% 与 74.61%, RebalFL 仍保持 88.30% 的较高模型准确率。值得注意的是, 在长尾分布预算设置下, MidFL 的模型准确率下降至 64.96%, PDPFL

为 73.17%, RebalFL 依旧取得 85.66% 的高模型准确率, 显著减缓了数据轻度异构带来的性能退化。

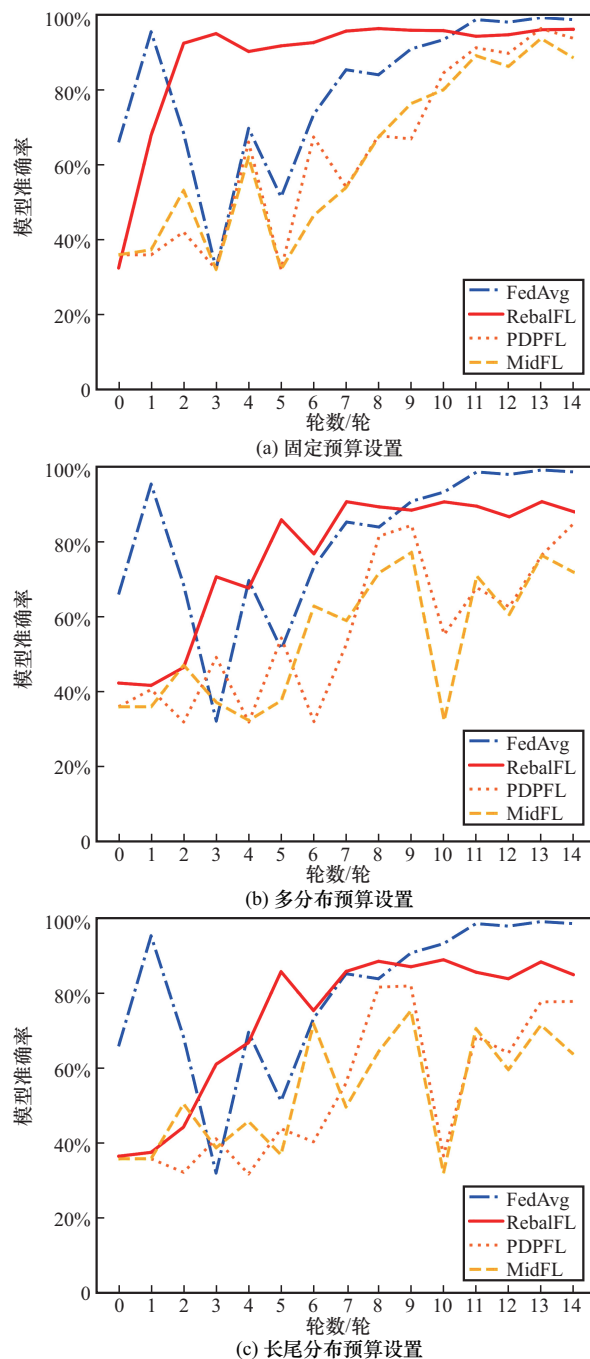


图5 Non-IID1 场景下各方法的模型准确率对比

表3

Non-IID1 场景下各方法的性能对比

隐私预算设置	FedAvg ^[42]		MidFL ^[43]		PDPFL ^[30]		RebalFL (本文方案)	
	模型准确率	运行时间/s	模型准确率	运行时间/s	模型准确率	运行时间/s	模型准确率	运行时间/s
固定预算	98.41%	235	89.25%	469	93.02%	610	95.36%	626
多分布预算	98.41%	235	69.48%	231	74.61%	268	88.30%	276
长尾分布预算	98.41%	235	64.96%	217	73.17%	243	85.66%	256

综合来看, RebalFL的优异表现源于两大设计的协同作用:首先, RebalFL采用结合抽样策略的PDP机制,通过保障低预算数据不会因隐私预算过早耗尽而退出训练,有效缓解了“数据被遗忘”的问题,确保了全局模型能获得更充分的数据贡献;其次, RebalFL设计的重平衡分簇策略,通过矫正客户端间的数据分布,直接抑制了模型漂移问题,使其在轻度异构场景下获得了更高的性能与更稳健的收敛。

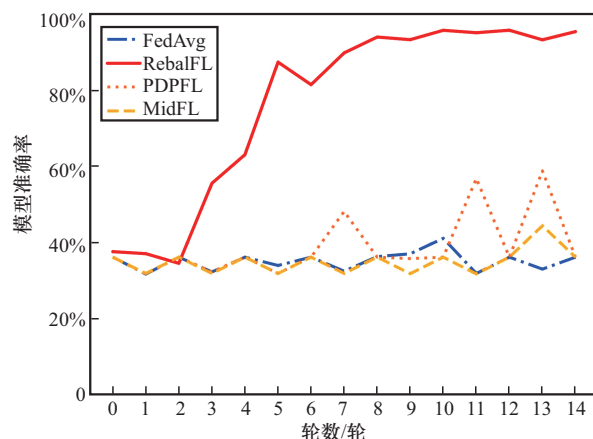
6.2.3 重度异构 Non-IID2 场景下的性能对比

在重度异构 Non-IID2 场景下, 次用户客户端间的数据分布差异被进一步放大, 各个次用户客户端的类别分布几乎互不重叠, 属于极端数据异构设置。实验结果如表4和图6所示, 其中, 本文方案的结果以加粗字体突出显示。传统FedAvg方法在此环境中因模型漂移问题出现明显性能退化, 导致模型准确率骤降至35.18%。在长尾分布预算中, MidFL模型准确率仅达到43.21%, PDPFL也下降至34.38%。相比之下, RebalFL在相同设置下表现出显著稳定性, 整体模型准确率始终大幅领先于基准方法, 在最极端的长尾分布预算设置下, RebalFL依旧能够维持82.07%的模型准确率。

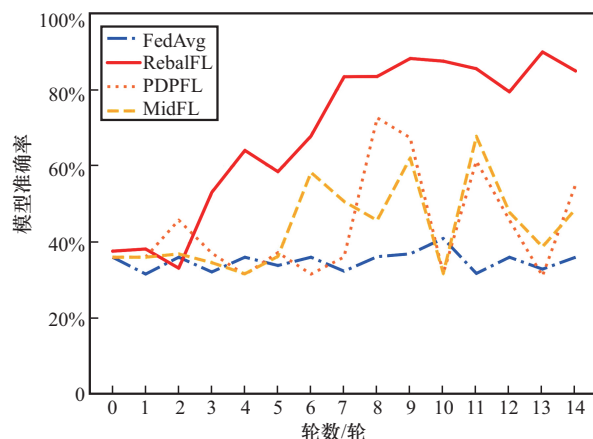
这一结果证明了RebalFL的客户端重平衡分簇策略能够有效识别并聚合具有互补分布特性的次用户客户端, 从源头上缓解模型漂移问题, 进而在多种 Non-IID 场景中能保证较高的模型准确率。

6.3 RebalFL缓解模型漂移有效性分析

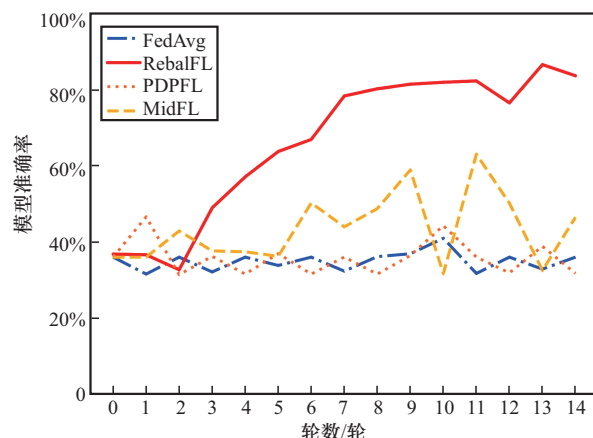
本节介绍RebalFL与PDPFL^[30]在模拟频谱数据二分类任务中的实验设置及结果分析。实验通过精确率、召回率、F1值、准确率及混淆矩阵结果进行分析评估, 重点考察RebalFL能否有效缓解因 Non-IID 和 PDP 机制引发的模型漂移问题。



(a) 固定预算设置



(b) 多分布预算设置



(c) 长尾分布预算设置

图6 Non-IID2场景下各方法的模型准确率对比

表4

Non-IID2场景下各方法的性能对比

隐私预算设置	FedAvg ^[42]		MidFL ^[43]		PDPFL ^[30]		RebalFL (本文方案)	
	模型准确率	运行时间/s	模型准确率	运行时间/s	模型准确率	运行时间/s	模型准确率	运行时间/s
固定预算	35.18%	223	38.95%	514	43.72%	642	94.42%	624
多分布预算	35.18%	223	45.24%	227	44.17%	256	84.73%	280
长尾分布预算	35.18%	223	43.21%	240	34.38%	280	82.07%	256

设置20个次用户客户端,模拟轻度异构场景。其中10个为混合数据客户端,各客户端包含70%的类别0(空闲状态)样本和30%的类别1(繁忙状态)样本;其余10个为单一数据客户端,其中5个仅包含类别0样本,另外5个仅包含类别1样本。

在模拟频谱二分类任务中,分别比较了RebalFL与PDPFL在精确率、召回率、F1值和准确率4个指标上的表现,结果如表5所示。表5显示,PDPFL尽管在类别0上保持了较高的精确率(100.00%),但模型整体性能不足,特别是对类别1样本的识别能力明显偏弱,其召回率仅为47.78%,导致整体F1值仅达到64.66%。相比之下,RebalFL在各项指标上均取得更优表现,整体准确率提升至88.13%,F1值提升至81.31%,展现出更均衡且更可靠的分类性能。可视化混淆矩阵如图7所示,可以更直观地看到RebalFL相较PDPFL方法有效提升了对类别1样本的正确识别能力。图7(a)显示,PDPFL将470个真实的类别1样本错误地判定为类别0样本,图7(b)展示了RebalFL能提升对类别1样本的正确识别能力,误判数量降至280个。

表5 模拟频谱二分类任务性能对比

方法	精确率	召回率	F1值	准确率
PDPFL ^[30]	100.00%	47.78%	64.66%	80.42%
RebalFL(本文方案)	99.20%	68.89%	81.31%	88.13%

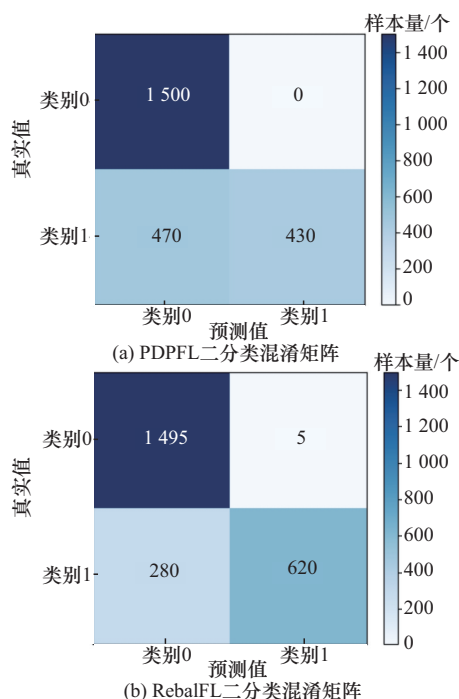


图7 PDPFL与RebalFL可视化混淆矩阵

以上实验结果充分表明,RebalFL通过重平衡分簇策略能够有效校正因隐私预算差异导致的模型漂移问题,在提升整体分类性能的同时,显著提高了对强隐私保护数据的识别效果。对于现实场景下的频谱感知任务,RebalFL在检测主用户繁忙状态时表现更可靠,从而能够在频段接入决策中更有效地避免对主用户通信的潜在干扰,具有更优的应用价值。

7 结束语

在频谱资源日益紧张背景下,FL为构建高效、隐私保护的协作频谱感知提供了可行路径。然而,如何在Non-IID数据与差异化隐私需求下,协同实现可靠的隐私保护与高精度的模型性能,是当前面临的核心挑战。基于此,本文提出了融合个性化差分隐私与重平衡分簇策略的联邦学习方案(RebalFL),引入记录级PDP机制,允许客户端依据数据敏感度灵活分配隐私预算,在提供严格隐私保护的同时,显著减少了过量噪声对模型效用的损害,设计了基于KL散度的客户端重平衡分簇策略,有效构建了分布均衡的训练簇以缓解模型漂移问题。实验结果表明,在Non-IID场景下,RebalFL取得了更好的隐私保护性能和模型可用性。

在未来的实际部署中,联邦学习系统可能面临拜占庭攻击、投毒攻击等恶意行为威胁。本文方案在设计中虽未专门针对恶意攻击进行加固,但其个性化差分隐私机制通过噪声注入可在一定程度上稀释恶意客户端上传的异常梯度,从而缓解部分投毒攻击的影响。此外,重平衡分簇策略基于客户端数据分布进行聚类,未来若结合客户端行为可信度评估,有望在簇内识别并隔离行为异常的节点,提升系统在分布攻击下的鲁棒性。后续工作将进一步探索DP在FL中的优化,并探索该框架在更广泛的无线通信场景,如频谱预测、资源分配等领域的应用,以进一步推动隐私保护的FL在频谱资源安全共享中的应用。

参考文献:

- [1] 王海涛, 茆习文, 张晨, 等. 基于天基干扰绘图的地空一体化系统频谱共享研究[J]. 通信学报, 2024, 45(3): 155-165.
Wang H T, Mao X W, Zhang C, et al. Study of spectrum sharing in integrated satellite-terrestrial system based on space-based interference cartography[J]. Journal on Communications, 2024, 45(3): 155-165.
- [2] Zhang W S, Wang Y, Chen X, et al. Collaborative learning-based spec-

- trum sensing under partial observations[J]. *IEEE Transactions on Cognitive Communications and Networking*, 2024, 10(5): 1843-1855.
- [3] Janu D, Singh K, Kumar S. Machine learning for cooperative spectrum sensing and sharing: a survey[J]. *Transactions on Emerging Telecommunications Technologies*, 2022, 33: e4352.
- [4] Hassan M U, Rehmani M H, Rehan M, et al. Differential privacy in cognitive radio networks: a comprehensive survey[J]. *Cognitive Computation*, 2022, 14(2): 475-510.
- [5] Bai L, Hu H B, Ye Q Q, et al. Membership inference attacks and defenses in federated learning: a survey[J]. *ACM Computing Surveys*, 2025, 57(4): 1-35.
- [6] Wasilewska M, Bogucka H, Poor H V. Secure federated learning for cognitive radio sensing[J]. *IEEE Communications Magazine*, 2023, 61(3): 68-73.
- [7] Yang H M, Ge M Y, Xue D Y, et al. Gradient leakage attacks in federated learning: research frontiers, taxonomy, and future directions[J]. *IEEE Network*, 2024, 38(2): 247-254.
- [8] Ouadrhiri A E, Abdelhadi A. Differential privacy for deep and federated learning: a survey[J]. *IEEE Access*, 2022, 10: 22359-22380.
- [9] Wu X, Zhang Y T, Shi M Y, et al. An adaptive federated learning scheme with differential privacy preserving[J]. *Future Generation Computer Systems*, 2022, 127: 362-372.
- [10] Bao T, Xu L, Zhu L H, et al. Successive point-of-interest recommendation with personalized local differential privacy[J]. *IEEE Transactions on Vehicular Technology*, 2021, 70(10): 10477-10488.
- [11] Costa L D S, Guimarães D A, Uchôa-Filho B F. On the signal-to-noise ratio wall of energy detection in spectrum sensing[J]. *IEEE Access*, 2022, 10: 16499-16511.
- [12] Arjoune Y, Kaabouch N. A comprehensive survey on spectrum sensing in cognitive radio networks: recent advances, new challenges, and future research directions[J]. *Sensors*, 2019, 19(1): 126.
- [13] Song Z H, Gao Y, Tafazolli R. A survey on spectrum sensing and learning technologies for 6G[J]. *IEICE Transactions on Communications*, 2021, 104(10): 1207-1216.
- [14] Charan C, Pandey R. Eigenvalue based double threshold spectrum sensing under noise uncertainty for cognitive radio[J]. *Optik*, 2016, 127(15): 5968-5975.
- [15] Nasser A, Hassan H A H, Chaaya J A, et al. Spectrum sensing for cognitive radio: recent advances and future challenge[J]. *Sensors*, 2021, 21(7): 2408.
- [16] Saber M, Rharras A E, Saadane R, et al. An optimized spectrum sensing implementation based on SVM, KNN and TREE algorithms[C]// *Proceedings of the 2019 15th International Conference on Signal-Image Technology & Internet-Based Systems (SITIS)*. Piscataway: IEEE Press, 2019: 383-389.
- [17] Liu C, Wang J, Liu X M, et al. Deep CM-CNN for spectrum sensing in cognitive radio[J]. *IEEE Journal on Selected Areas in Communications*, 2019, 37(10): 2306-2321.
- [18] Solanki S, Dehalwar V, Choudhary J, et al. Spectrum sensing in cognitive radio using CNN-RNN and transfer learning[J]. *IEEE Access*, 2022, 10: 113482-113492.
- [19] Li H N, Hu T H, Chen J X, et al. Privacy preserving algorithm for spectrum sensing in cognitive vehicle networks[J]. *Chinese Journal of Electronics*, 2024, 33(1): 30-42.
- [20] Kairouz P, McMahan H B. Advances and open problems in federated learning[J]. *Foundations and Trends in Machine Learning*, 2021, 14(1/2): 1-210.
- [21] Chen Z B, Xu Y Q, Wang H B, et al. Federated learning-based cooperative spectrum sensing in cognitive radio[J]. *IEEE Communications Letters*, 2022, 26(2): 330-334.
- [22] Song Y F, Chang H H, Liu L J. Federated dynamic spectrum access through multi-agent deep reinforcement learning[C]// *Proceedings of the GLOBECOM 2022-2022 IEEE Global Communications Conference*. Piscataway: IEEE Press, 2022: 3466-3471.
- [23] Shi Y, Sagduyu Y E, Erpek T. Federated learning for distributed spectrum sensing in NextG communication networks[PP]. V1. (2022-04-06) [2025-12-01]. arXiv: arXiv. 2204.03027.
- [24] Li T, Sahu A K, Zaheer M, et al. Federated optimization in heterogeneous networks[J]. *Proceedings of Machine Learning and Systems*, 2020, 2: 429-450.
- [25] Wang J Y, Liu Q H, Liang H, et al. Tackling the objective inconsistency problem in heterogeneous federated optimization[J]. *Advances in Neural Information Processing Systems*, 2020, 33: 7611-7623.
- [26] 李梦倩, 田有亮, 张军鹏, 等. 基于零集中差分隐私的联邦学习激励方案[J]. *通信学报*, 2025, 46(1): 79-92.
- Li M Q, Tian Y L, Zhang J P, et al. Incentive scheme for federated learning based on zero-concentrated differential privacy[J]. *Journal on Communications*, 2025, 46(1): 79-92.
- [27] Fu J, Hong Y, Ling X P, et al. Differentially private federated learning: a systematic review[PP]. V4. (2025-10-02) [2025-12-01]. arXiv: arXiv. 2405.08299.
- [28] Fu J, Ye Q Q, Hu H B, et al. DPSUR: accelerating differentially private stochastic gradient descent using selective update and release[PP]. V2. (2023-11-29) [2025-12-01]. arXiv: arXiv. 2311.14056.
- [29] Liu J X, Lou J, Xiong L, et al. Projected federated averaging with heterogeneous differential privacy[J]. *Proceedings of the VLDB Endowment*, 2021, 15(4): 828-840.
- [30] Liu J X, Lou J, Xiong L, et al. Cross-silo federated learning with record-level personalized differential privacy[C]// *Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security*. New York: ACM Press, 2024: 303-317.
- [31] Jorgensen Z, Yu T, Cormode G. Conservative or liberal? Personalized differential privacy[C]// *Proceedings of the 2015 IEEE 31st International Conference on Data Engineering*. Piscataway: IEEE Press, 2015: 1023-1034.
- [32] Boenisch F, Dziedzic A, Mühl C, et al. Have it your way: individualized privacy assignment for DP-SGD[C]// *Proceedings of the Advances in Neural Information Processing Systems 36*. Massachusetts: MIT Press, 2023: 19073-19103.
- [33] Aygöl M A, Çırpan H A, Arslan H. Machine learning-based spectrum occupancy prediction: a comprehensive survey[J]. *Frontiers in Communications and Networks*, 2025, 6: 1482698.
- [34] Zhao Y, Chen J J. A survey on differential privacy for unstructured data content[J]. *ACM Computing Surveys*, 2022, 54(10s): 1-28.
- [35] Niu B, Chen Y H, Wang B Y, et al. AdaPDP: adaptive personalized differential privacy[C]// *Proceedings of the IEEE INFOCOM 2021-IEEE Conference on Computer Communications*. Piscataway: IEEE Press, 2021: 1-10.
- [36] Liu Z X, Li T, Smith V, et al. Enhancing the privacy of federated learning with sketching[PP]. V1. (2019-11-05) [2025-12-01]. arXiv: arXiv. 1911.01812.

- [37] Hayes J, Melis L, Danezis G, et al. LOGAN: membership inference attacks against generative models[J]. *Proceedings on Privacy Enhancing Technologies*, 2019, 2019(1): 133-152.
- [38] Wang Z B, Song M K, Zhang Z F, et al. Beyond inferring class representatives: user-level privacy leakage from federated learning[C]//*Proceedings of the IEEE Conference on Computer Communications*. Piscataway: IEEE Press, 2019: 2512-2520.
- [39] Mironov I, Talwar K, Zhang L. Rényi differential privacy of the sampled Gaussian mechanism[PP]. V1. (2019-08-28) [2025-12-01]. arXiv: arXiv.1908.10530.
- [40] Mironov I. Rényi differential privacy[C]//*Proceedings of the 2017 IEEE 30th Computer Security Foundations Symposium (CSF)*. Piscataway: IEEE Press, 2017: 263-275.
- [41] Arjouni Y, Kaabouch N. On spectrum sensing, a machine learning method for cognitive radio systems[C]//*Proceedings of the 2019 IEEE International Conference on Electro Information Technology (EIT)*. Piscataway: IEEE Press, 2019: 333-338.
- [42] McMahan B, Moore E, Ramage D, et al. Communication-efficient learning of deep networks from decentralized data[C]//*Artificial Intelligence and Statistics*. New York: PMLR, 2017: 1273-1282.
- [43] Geyer R C, Klein T, Nabi M. Differentially private federated learning: a client level perspective[PP]. V2. (2018-03-01) [2025-12-01]. arXiv: arXiv.1712.07557.

[作者简介]



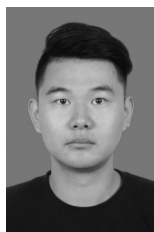
唐湘云 (1994-), 女, 湖南永州人, 博士, 中央民族大学副教授, 主要研究方向为人工智能安全、数据安全、隐私保护等。



康嘉文 (1992-), 男, 广东茂名人, 博士, 广东工业大学教授, 主要研究方向为人工智能、信息安全等。



韩旭 (2002-), 女, 广西河池人, 中央民族大学硕士生, 主要研究方向为联邦学习、隐私保护等。



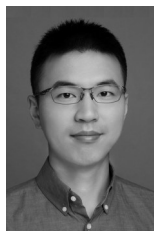
张焘 (1994-), 男, 安徽马鞍山人, 博士, 北京交通大学副教授, 主要研究方向为区块链、联邦学习等。



刘寅秋 (1998-), 男, 江苏徐州人, 博士, 新加坡南洋理工大学副研究员, 主要研究方向为区块链安全、边缘智能等。



孙庚 (1990-), 男, 吉林长春人, 博士, 吉林大学教授, 主要研究方向为低空无线网络、移动边缘计算等。



焦雨涛 (1992-), 男, 江苏南京人, 博士, 中国人民解放军陆军工程大学副教授, 主要研究方向为智能短波通信、电磁频谱感知、无线联邦学习等。