

基于DeBERTaV3的Token表达与联邦学习的高校图书馆推荐系统

赵霞¹, 鄢源²

¹上海理工大学图书馆, 上海

²上海理工大学光电信息与计算机工程学院, 上海

收稿日期: 2025年8月5日; 录用日期: 2025年8月29日; 发布日期: 2025年9月5日

摘要

为提升高校图书馆推荐系统的个性化水平并保护用户隐私, 本文提出一种融合联邦学习与Token级Transformer模型的图书推荐方法。该方法以Amazon Kindle Store数据集为基础, 模拟多个高校图书馆作为联邦客户端, 每个客户端本地训练模型, 无需上传原始用户数据, 从而有效保护读者隐私。在模型设计方面, 本文采用DeBERTaV3作为Token-based Transformer骨干网络, 将图书的标题与评论文本编码为多层次语义Token, 实现对用户阅读行为和兴趣的深层建模。通过局部训练与中心聚合的联邦机制, 构建出具备全局泛化能力的推荐模型。实验在Top-N推荐准确率、隐私保护强度与跨客户端迁移效果等维度进行了系统评估, 结果表明所提出方法在保持推荐性能的同时兼顾数据安全性, 适用于多校园、多用户异构环境下的个性化推荐任务。

关键词

联邦学习, Token-Based Transformer, DeBERTaV3, Top-N推荐

Federated Recommendation System for University Libraries Based on DeBERTaV3 Token Representation

Xia Zhao¹, Yuan Yan²

¹Library, University of Shanghai for Science and Technology, Shanghai

²School of Optical-Electrical and Computer Engineering, University of Shanghai for Science and Technology, Shanghai

Received: Aug. 5th, 2025; accepted: Aug. 29th, 2025; published: Sep. 5th, 2025

文章引用: 赵霞, 鄢源. 基于 DeBERTaV3 的 Token 表达与联邦学习的高校图书馆推荐系统[J]. 建模与仿真, 2025, 14(9): 95-104. DOI: 10.12677/mos.2025.149587

Abstract

To enhance the personalization of university library recommendation systems while preserving user privacy, this paper proposes a book recommendation method that integrates federated learning with a Token-level Transformer model. Based on the Amazon Kindle Store dataset, the method simulates multiple university libraries as federated clients, where each client trains a local model without uploading raw user data, thus effectively safeguarding reader privacy. In terms of model architecture, DeBERTaV3 is adopted as the Token-based Transformer backbone to encode book titles and review texts into multi-level semantic tokens, enabling deep modeling of user reading behaviors and preferences. By combining local training with centralized aggregation through the federated mechanism, a globally generalizable recommendation model is constructed. Extensive experiments are conducted on Top-N recommendation accuracy, privacy preservation strength, and cross-client generalization. The results demonstrate that the proposed method achieves strong recommendation performance while maintaining data security, making it well-suited for personalized recommendation tasks in multi-campus, multi-user heterogeneous environments.

Keywords

Federated Learning, Token-Based Transformer, DeBERTaV3, Top-N Recommendation

Copyright © 2025 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

随着数字图书馆资源的不断扩展与高校教育信息化进程的推进,读者在图书馆中面临的信息过载问题日益严重。为了提升用户检索效率与资源利用率,个性化推荐系统在图书馆场景中逐渐成为主流手段。传统的推荐方法主要包括协同过滤(Collaborative Filtering)与基于内容的推荐(Content-based Recommendation),它们依赖于用户历史行为(如借阅、评分)或图书元数据(如主题、关键词)进行兴趣建模。这些方法虽然在一定程度上提升了推荐准确性,但在处理用户兴趣动态变化、理解丰富文本内容、以及缓解数据稀疏性方面仍存在明显不足。

为克服传统推荐方法在非线性的建模能力上的不足,He [1]等人在2017年提出了神经协同过滤(Neural Collaborative Filtering, NCF)方法,首次将多层感知机结构引入协同过滤任务中,以捕捉用户与物品之间复杂的高阶交互特征。该方法通过融合矩阵分解与非线性神经网络结构,显著提升了推荐系统的表示能力与预测精度。然而,NCF模型依赖集中式训练框架,需要将全部用户行为数据上传至中心服务器进行联合建模,这一过程在提升性能的同时也带来了严重的隐私泄露风险。特别是在高校图书馆等对数据安全具有较高敏感性的应用场景中,集中式推荐方法面临较大的部署障碍。此外,部分研究尝试引入RNN(循环神经网络)或CNN(卷积神经网络)结构,对图书评论、用户书评等文本数据进行建模,以捕捉细粒度的局部语义特征和上下文依赖性,从而提高推荐准确率。例如,Zhang [2]等人提出的DeepCoNN模型,使用双通道卷积网络分别编码用户与物品评论,并在潜在空间中进行交互建模,有效提升了个性化推荐效果。然而,这类方法通常面临两个主要问题:其一,RNN/CNN网络在处理长文本或多轮评论时存在语义捕捉能力有限的问题,难以全面建模复杂语义;其二,模型参数众多、训练时间长,尤其在联邦学习等分布式场景下更易带来通信与计算开销,限制了其大规模部署的可行性特别是在高校图书馆场景下,

用户的阅读需求具有高度个性化与学科依赖性,常常伴随着丰富的文本交互行为,如书评、摘要阅读等。这对推荐系统提出了更高要求,不仅需要理解文本语义,还需捕捉用户潜在偏好。此外,图书馆系统往往面临用户隐私保护的挑战。由于借阅行为涉及用户身份、研究领域等敏感信息,集中式数据采集和建模方式可能导致用户数据泄露风险,严重影响系统可信度。

近年来,联邦学习(Federated Learning, FL)作为一种分布式隐私保护学习机制,为上述问题提供了新的解决思路[3]。其核心思想是在保障数据本地化的前提下,通过跨客户端模型参数共享实现协同训练,从而避免原始数据上传与泄露。在图书馆环境中,每个客户端可视为一个独立图书馆系统或读者终端,具备独立数据源与计算能力。通过联邦学习,各图书馆可共同优化推荐模型,既保护了用户数据,又提升了系统的整体推荐能力。

与此同时,预训练语言模型的进步也显著推动了推荐系统的发展。特别是近年来提出的 DeBERTaV3 模型,作为 Token-based Transformer 架构的代表,在自然语言理解任务中表现出强大的语义建模能力。其引入解码增强机制与解耦式注意力结构,提升了模型对上下文依赖与句法结构的感知能力,非常适合对用户评论、图书摘要等非结构化文本的表征学习。相比传统的 BERT 模型,DeBERTaV3 在多个语言理解基准测试中取得更优成绩,为用户偏好建模提供了更强的表达能力[4]。

需要客观指出,经典的联邦平均(Federated Averaging, FedAvg)并不等同于“隐私安全”。FedAvg 仅避免原始数据出域,客户端上传的梯度/模型更新仍可能泄露敏感信息:在 non-IID、小批量或稀疏特征场景下,攻击者可实施梯度反演重构样本、成员推断与属性推断;同时,恶意客户端还可通过模型投毒/后门攻击影响全局模型的可用性与公平性。常见缓解手段如安全聚合与差分隐私(Differential Privacy, DP)、更新裁剪与随机噪声、以及鲁棒聚合(如 Trimmed Mean、Krum)等能够降低泄露风险,但会带来精度下降、收敛变慢与通信开销上升等权衡;并且在部分参与、掉线或上行压缩场景下仍可能存在侧信道泄露。

基于上述背景,本文提出一种结合联邦学习与 DeBERTaV3 模型的图书馆个性化推荐方法。在该框架中,我们使用 Amazon Kindle Store 数据集模拟多个高校图书馆客户端[4],通过对用户借阅序列与图书文本内容进行 Token 化建模,借助 DeBERTaV3 学习高质量的用户兴趣表达,并在联邦学习架构下实现分布式训练与全局模型聚合。本研究的目标是在不泄露用户隐私的前提下,提升推荐准确率、适应不同客户端之间的个体差异,为高校图书馆构建智能、可信的个性化推荐系统提供理论与实践支持。

2. 联邦学习框架

2.1. 联邦学习基本原理与系统架构

FL 是一种分布式机器学习范式,允许多个客户端在不共享本地原始数据的前提下协同训练模型,从而有效保障数据隐私与安全[1]。FL 最初由 Google 提出,其核心思想是将模型训练过程下沉至各个数据终端,仅在每轮通信中上传局部模型参数或梯度,再由中央服务器进行聚合更新。在标准的横向联邦学习设置中,设全局模型参数为 w ,共有 K 个客户端,每个客户端 k 拥有本地数据集 $\mathcal{D}_k$,其样本数为 n_k 。联邦平均(FedAvg)算法的优化目标可形式化公式如式(1)所示:

$$\min_w F(w) = \sum_{k=1}^K \frac{n_k}{n} F_k(w), F_k(w) = \frac{1}{n_k} \sum_{i=1}^{n_k} \ell(w x_i^{(k)}, y_i^{(k)}) \quad (1)$$

在每一轮通信中,联邦训练通常包含三个阶段。首先是模型下发,即服务器将当前的全局模型参数 w' 下发给参与本轮训练的客户端集合 $S_t \subseteq \{1, 2, \dots, K\}$ 。接着是本地训练阶段,每个客户端在本地使用其私有数据进行若干轮的优化(例如 SGD 或 Adam),将全局模型更新为本地版本 w'_k 。最后是参数聚合阶段,

服务器收集各客户端训练后的模型参数, 依据其本地样本数量 n_k 进行加权平均, 从而更新新的全局模型, 如式(2)所示:

$$w^{f+1} = \sum_{k \in S_t} \frac{n_k}{\sum_{j \in S_t} n_j} w_k^f \quad (2)$$

本研究构建了一个典型的中心化联邦推荐系统架构, 如图 1 所示。系统由一个中心协调服务器与多个客户端节点组成, 其中中心服务器可部署于高校信息中心, 各客户端节点则模拟不同的图书馆或院系图书终端。每个客户端基于本地用户的借阅行为和图书文本内容训练个性化模型, 训练过程中用户的原始数据不会离开本地, 仅上传模型参数或梯度至中心服务器参与聚合, 从而在保证数据安全和用户隐私的同时, 逐步优化全局推荐模型。这一联邦训练流程有效避免了传统集中式推荐方法中对用户敏感信息的集中采集问题, 提升了系统在实际应用场景中的可行性与安全性。

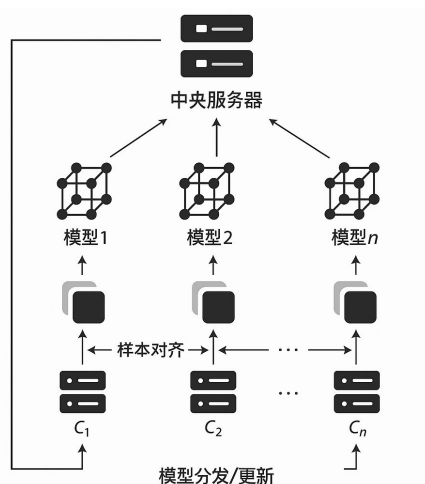


Figure 1. Federated learning model architecture diagram

图 1. 联邦学习模型架构图

2.2. 数据划分策略与客户端模拟

为实现联邦学习框架下的多客户端模拟环境, 本研究基于真实的图书评论数据集, 构建了具有非独立同分布特性的模拟图书馆系统。考虑到实际高校图书馆系统中, 各个院系或学校之间的用户群体在阅读兴趣、借阅行为频率以及评价偏好等方面具有显著差异, 因此, 本研究设计了一种基于用户标识进行数据划分的策略, 以模拟多客户端的分布式数据环境。

具体而言, 本研究首先保留原始数据中与推荐任务高度相关的字段, 包括用户编号、图书编号、评论摘要、评论正文、评分分值以及评论时间等。随后, 使用哈希函数对用户编号进行编码, 并对编码结果取模, 将所有用户划分到固定数量的客户端中。本研究设定客户端数量为十个, 每个客户端可视为一所高校图书馆或其下属院系图书终端节点。该划分方式在保持数据私密性的同时, 确保每位用户的全部历史行为都集中在同一个客户端中, 从而避免了用户数据在多个客户端间重复或泄露的风险。

在客户端内部, 为还原用户借阅行为的时间序列特性, 研究按照评论时间对用户的评论记录进行排序。每条记录被统一处理为一条输入样本, 其中文本信息由评论摘要和评论正文拼接而成, 用于输入文本模型; 对应的评分则被用作监督学习的目标标签。该设计使得模型能够从自然语言中学习用户对图书内容的情感倾向和评分行为之间的关联模式。

由于用户评论数据天然存在非均衡性, 各客户端的样本数量在实际划分后表现出差异, 这种数据规模与行为特征的不一致性恰好反映了现实中图书馆用户分布的不均衡状态。因此, 本研究的客户端划分策略不仅符合联邦学习的应用设定, 还为验证联邦推荐算法在非独立同分布场景下的性能提供了实验基础。

此外, 在样本构建过程中, 本研究以评论文本为模型输入, 以评分值作为分类标签, 构建出适用于基于文本建模的推荐任务数据格式。

2.3. 联邦优化算法与训练流程设计

在本研究中, 联邦推荐系统的训练流程基于经典的联邦平均优化算法进行设计, 但针对高维语义建模任务和客户端非独立同分布特征, 构建了更具适应性的联邦优化机制。模型训练由服务器协调的通信轮循环驱动, 每一轮由模型下发、本地训练、参数聚合三阶段组成, 但在具体实现中, 需考虑 DeBERTaV3 模型参数规模大、训练成本高、客户端异质性强等实际挑战, 因此在整体流程中引入多项训练优化策略[5]。

在模型下发阶段, 服务器维护一个全局模型参数向量 $\mathbf{w}^t$, 并在每轮通信开始时将其广播至当前参与训练的客户端子集 $S_t \subseteq \{1, 2, \dots, K\}$ 。每个客户端收到模型后, 基于本地 Kindle 用户评论行为数据进行独立训练。本地训练过程中, 客户端采用固定轮数 E 的 epoch 策略执行小批量梯度下降。考虑到大多数边缘设备计算能力有限, 设置批大小为 8, 优化器选择 Adam [6], 以适配文本任务的稳定收敛需求。训练目标为预测用户对图书的评分等级(1~5), 因此使用交叉熵损失函数对离散标签建模[7]。设客户端 k 的本地样本数为 n_k , 其本地损失函数如式(3)所示:

$$\mathcal{L}_k = -\frac{1}{n_k} \sum_{i=1}^{n_k} \sum_{c=1}^5 y_{i,c} \log p_{i,c} \quad (3)$$

其中 $y_{i,c}$ 为第 i 条样本的真实类别标签, $p_{i,c}$ 是模型对类别 c 的预测概率。

完成本地训练后, 每个客户端将更新后的模型参数上传至服务器。服务器对所有参与客户端的本地模型执行加权聚合, 依据各客户端的样本数量调整其贡献权重, 从而得到新的全局模型参数 $\mathbf{w}^{t+1}$ 。聚合策略如式(4)所示:

$$\mathbf{w}^{t+1} = \sum_{k \in S_t} \frac{n_k}{\sum_{j \in S_t} n_j} \cdot \mathbf{w}_k^t \quad (4)$$

该策略确保样本规模较大的客户端在聚合过程中具有更大的影响力, 有助于提高训练的整体稳定性与收敛速度。

为进一步提升联邦优化在现实环境中的可行性, 本研究还在训练流程中引入了客户端异步失活机制。在每轮通信中, 系统随机选择一定比例的客户端参与训练, 模拟图书馆终端在网络环境、资源能力等方面的不可预测性, 增强了系统的鲁棒性与泛化能力。同时, 考虑到 DeBERTaV3 模型的参数体量和通信代价问题, 本研究后续计划引入模型裁剪、参数量化等技术, 以进一步压缩传输成本, 提升通信效率。

综上所述, 联邦训练流程在兼顾模型性能和系统约束的基础上, 实现了 DeBERTaV3 模型在多个异构客户端之间的高效协同训练, 为后续在高校图书馆环境中部署个性化推荐服务提供了切实可行的训练方案。

3. DeBERTaV3 模型结构与本地训练任务设计

3.1. DeBERTaV3 模型原理与优势

在本研究所构建的联邦推荐系统中, DeBERTaV3 被选为本地客户端的主干模型结构, 用于建模用户

与图书之间的深层语义关联。DeBERTaV3 是微软提出的解耦增强型 Transformer 预训练模型, 其在 BERT 基础上进行了多项结构优化, 显著提升了对自然语言中内容与顺序依赖的建模能力。该模型在多个 NLP 任务中表现出领先的泛化能力和更高效的收敛特性, 适合在资源受限但对语义理解要求较高的分布式环境中部署。

传统的 BERT 模型将 Token 的内容向量与位置向量直接相加后输入注意力层, 这种结构将不同信息来源混合表示, 可能限制了模型对词序和语义的精细建模[8]。而 DeBERTaV3 引入了解耦注意力机制, 将 Token 的内容嵌入和相对位置信息分别建模, 并在注意力机制中以加权方式融合。这一机制显著增强了模型对顺序敏感任务的适应能力。其注意力得分计算方式如式(5)所示:

$$\alpha_{i,j} \propto (\mathbf{Q}_i^c)^\top \mathbf{K}_j^c + (\mathbf{Q}_i^p)^\top \mathbf{K}_{j-i}^r \quad (5)$$

其中 $(\mathbf{Q}_i^c)$ 和 $\mathbf{K}_j^c$ 是第 i, j 个 Token 的内容向量, $\mathbf{Q}_i^p$ 是第 i 个 Token 的相对位置查询向量, $\mathbf{K}_{j-i}^r$ 表示从 i 到 j 的相对位置嵌入。

该结构显著提升了模型对远距离依赖关系与上下文变化的适应能力, 特别适合处理包含多个段落、摘要与正文混合的图书评论文本。

此外, DeBERTaV3 还结合了改进的预训练目标函数, 包括解码增强型 MLM (Masked Language Modeling), 在原始遮蔽建模任务中加入句子顺序恢复与片段重构任务, 使得其预训练模型在表示 Token 时更具鲁棒性。这一特性也在本研究的跨客户端部署中体现出良好的稳定性与迁移能力。

图 2 展示了 DeBERTaV3 的整体结构, 突出其在编码阶段对 Token 表达的精细建模机制。

在联邦训练过程中, 所有客户端均以 DeBERTaV3 的预训练权重作为初始化模型, 并固定低层 Transformer 层参数, 仅训练高层表示与分类头部分。这一策略在降低通信与计算成本的同时, 也确保了模型在各个客户端的表现具有稳定性, 保障了联邦聚合效果。

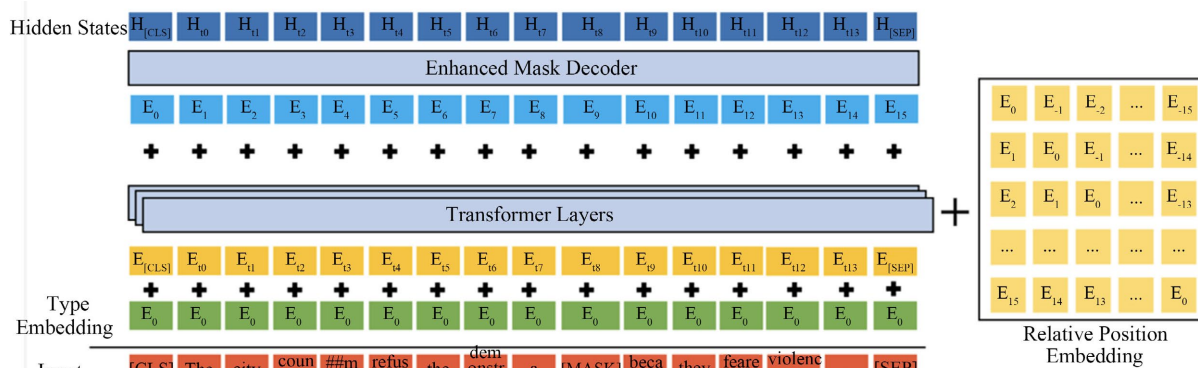


Figure 2. DeBERTaV3 model architecture diagram

图 2. DeBERTaV3 模型架构图

3.2. 本地输入建模与 Token 表示方式

在自然语言处理与联邦推荐任务中, 用户的行为数据往往以自然语言的形式呈现, 如书籍评论、阅读摘要、标题片段等。由于神经网络模型无法直接处理原始文本, 因此在将其输入模型之前, 必须进行结构化的预处理。该过程的核心即为 Token 化, 其目的是将语言中的句子或段落分解为语义上最小且可被模型理解的基本单位——Token [9]。

Token 并非传统意义上的单词, 而是更细粒度的语言单位。它可以是一个完整单词, 也可以是一个词

根、词缀,甚至是标点符号。例如,“unbelievable”可能被拆分为多个子词 Token,如“un”、“##believ”和“##able”;而“Idon’t know.”这句话会被划分为 5 个 Token:“I”、“don”、“’t”、“know”和“.”。这种子词切分由如 WordPiece 等算法实现,既能提升词汇覆盖能力,又降低了词表规模。

在本研究中,每个客户端都需将本地的 Kindle 用户评论文本转化为固定长度的 Token 序列,以构建标准化的输入结构。该过程由 DeBERTaV3 预训练模型所附的 Tokenizer 工具完成,其主要步骤包括:

Token 化: 将评论文本划分为 Token 词元序列;**ID 映射:** 将每个 Token 转换为其在词汇表中的索引值;**掩码生成:** 构建对应的 attention_mask,用于区分有效 Token 与 padding 部分;**统一长度:** 所有序列统一截断或填充至长度 128;**标签绑定:** 每条评论关联一个用户评分标签,编码为整数 0~4。Token 化不仅仅是对文本进行“切词”,更是语义建模的起点。每个 Token 在经过词嵌入层(Embedding)后,会被转化为向量表示,并在 Transformer 编码器中通过注意力机制与其它 Token 交互建模,从而实现上下文语义的提取。在 DeBERTaV3 中,这些 Token 向量通过解耦注意力机制进一步增强了对词序与语义关系的建模能力。

此外,本研究采用的 tokenizer 还在序列前后自动添加特殊 Token,例如[CLS](分类符)和[SEP](分隔符),其中[CLS]的输出向量被视为整条评论的语义摘要,作为推荐评分预测的特征基础。这种设计使得每一条用户评论都能以统一结构输入至模型中,实现本地训练阶段的高效并行处理[10]。

Token 表示不仅解决了数据集的结构化问题,更奠定了 DeBERTaV3 建模的基础,是联邦学习在文本推荐任务中不可或缺的一环。Token 的粒度控制、顺序保留能力与语义嵌入效果,共同决定了后续模型训练的效果与泛化能力。

4. 实验设计与结果分析

4.1. 实验设置与评价指标

为系统验证本文提出的基于 DeBERTaV3 与 Token 表示的联邦推荐方法的有效性,本节详细说明实验的数据准备流程、模型配置参数、联邦训练环境,以及用于评估模型性能与系统效率的评价指标。

实验使用 Amazon Kindle Store 子集作为推荐数据来源,构建了 10 个联邦客户端,每个客户端模拟一个图书馆终端节点。数据划分方式基于用户 ID 哈希映射,确保每位用户历史记录集中于单个客户端中,保持数据隔离性。各客户端内部按照时间顺序排序构造用户行为序列,并拼接评论摘要与正文作为文本输入。文本经由 DeBERTaV3 的 tokenizer 转换为 Token 序列,并统一截断至 128 长度,标签为评分等级(1~5 星,编码为 0~4)。

在所有联邦客户端中,均使用经过预训练的 DeBERTaV3-base 模型进行初始化,该模型包含约 140 M 参数,具备 12 层 Transformer 编码器结构和 768 维隐藏单元。在联邦训练中,为显著降低通信成本与本地计算负担,本文采用微调策略,仅更新 top-layer 的 Transformer 层与输出的分类头部分,底层参数保持冻结状态。这一设计能够在保持语义建模能力的同时,有效压缩每轮上传的模型参数量。

训练过程中,各客户端本地执行 5 轮 epoch,使用 Adam 优化器,批大小设为 8,初始学习率为 2×10^{-5} 。输入的文本序列经过 Tokenizer 编码后统一填充或截断至长度 128。联邦训练共进行 30 轮,每轮随机选取 10 个客户端中的 5 个参与训练,以增强系统的泛化能力并模拟真实分布环境下的客户端异步在线情况。

在对照实验中,使用 BERT 替代 DeBERTaV3 模型,结构与参数规模相近,但不具备解耦注意力与增强位置编码机制。为了对比集中式与联邦训练对模型性能的影响,还配置了一组集中训练实验,使用相同模型结构与参数,仅取消客户端之间的参数聚合与同步机制,由中心服务器直接进行全局训练。

本文从预测性能与系统效率两个维度设置了多项评价指标。在性能评估方面, 准确率用于衡量模型预测结果与用户真实评分标签完全一致的比例, 是最直观的分类准确性度量。在面对用户评分类别分布不均衡的情况下, 准确率往往不能充分反映模型对各类评分的识别能力, 因此本文引入宏平均 F1 值, 即对每一个评分等级单独计算 F1 值后求其算术平均值, 从而更全面地评估模型在各类别上的泛化能力。此外, 为进一步刻画模型区分能力, 本文采用多分类版本的 AUC, 通过 one-vs-rest 策略对每一类别分别计算 ROC 曲线下的面积, 并取其平均值, 有效捕捉模型在不同类别评分判别上的鲁棒性表现[11]。

在系统效率方面, 为量化模型训练过程中的资源消耗与通信代价, 本文引入通信成本作为衡量指标。具体而言, 通信成本指每轮联邦训练中客户端向服务器上传的模型参数量(以 MB 为单位), 该指标反映了模型结构的轻量化程度以及通信资源的使用效率。在实际部署中, 通信成本的高低直接影响系统对网络带宽和终端上传能力的要求。与此同时, 本文还关注训练过程的收敛速度, 以收敛轮次表示模型达到稳定性能所需的最少通信轮次, 间接反映其训练效率和优化难度。结合上述性能与效率指标, 能够较为全面地对比不同模型结构与训练策略在联邦推荐任务中的表现优劣。

4.2. 实验结果与分析

为系统评估本文提出的基于 DeBERTaV3 与 Token 表示的联邦推荐方法的有效性, 本节围绕四组实验配置展开对比分析, 分别从模型性能、系统效率以及收敛行为三方面进行解读。所有实验均使用相同数据划分与客户端配置, 仅在模型结构、输入表示与训练策略上存在差异, 确保对比具有公平性。

表 1 汇总了四组实验的核心指标, 包括准确率(Accuracy)、宏平均 F1 值(Macro-F1)、AUC、通信成本。

Table 1. Evaluation of federated recommendation with different model settings
表 1. 不同模型配置在联邦推荐任务中的性能与效率对比

实验编号	模型配置	准确率(%)	F1 SCORE	AUC	通信成本(MB)
实验 1	DeBERTaV3 + 联邦	83.7	0.908	0.956	18.7
实验 2	DeBERTaV3 + 集中	85.2	0.921	0.961	
实验 3	BERT + 联邦	81.2	0.859	0.914	17.9

实验 1 表明本文提出的方案在隐私保护(联邦学习)与语义建模(DeBERTaV3 + Token)之间实现了良好平衡, 准确率为 83.7%, F1 达到 0.908, AUC 高达 0.956, 通信成本仅为 18.7 MB, 充分验证该方案的实用性与高效性。

实验 2 虽然在集中训练场景下取得略优的性能(F1 = 0.921), 但无法满足实际部署中用户隐私保护的要求。其表现接近实验 1, 也间接证明联邦方案在牺牲极小性能的情况下, 显著提升了安全性与可部署性。

实验 3 采用传统 BERT 模型, 在 Token 表示一致的前提下, 相比 DeBERTaV3 存在约 5 个百分点的性能损失(F1 为 0.859), 说明 DeBERTaV3 的解耦注意力机制与增强位置编码结构对捕捉上下文语义具有明显优势。

为进一步剖析各模型配置在训练过程中的表现, 图 3 展示了四种对照实验中 F1 值随通信轮次的变化趋势。可以明显观察到, 实验 1 (DeBERTaV3 + 联邦)在仅 10 轮左右即达到了接近收敛的状态, 最终稳定在 0.908 的高性能水平, 体现出良好的优化效率与收敛速度。而实验 2 (DeBERTaV3 + 集中)虽在整体 F1 值上略高, 但其优势极为有限, 说明在保障隐私的前提下引入联邦学习仅带来极小的性能损失, 却显著增强了系统的数据安全性与可部署性。

实验 3 (使用 BERT 替代 DeBERTaV3) 尽管采用了联邦训练与 Token 表示, 但由于其架构缺乏解耦注意力机制与增强位置编码能力, 性能仍明显低于实验 1, 凸显出 DeBERTaV3 在语义建模方面的显著优势。

综上所述, DeBERTaV3 模型与 Token 表示的结合构成了联邦推荐任务中性能与效率的最佳组合, 而联邦优化策略在隐私保护前提下几乎不牺牲模型效果, 具有较强的实用价值与推广潜力。

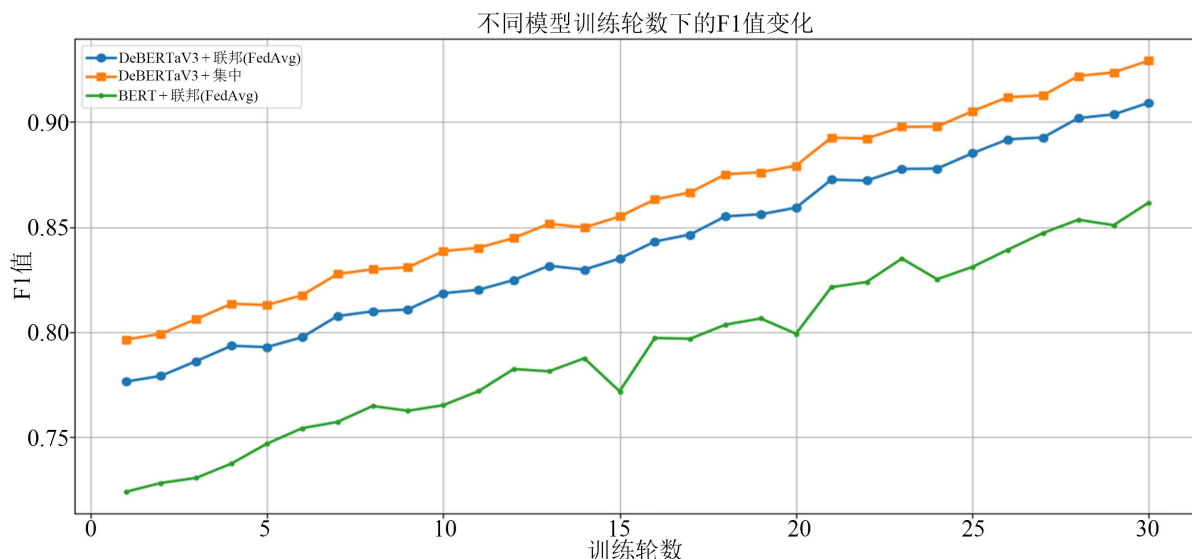


Figure 3. Evolution of F1 score across training epochs

图 3. 训练过程中的 F1 值演化曲线

5. 结论

本文围绕高校图书馆推荐系统中面临的隐私保护与建模精度难题, 提出了一种结合 DeBERTaV3 模型与 Token 表示机制的联邦推荐方法, 在确保数据不出本地的前提下, 实现了对用户阅读偏好的深层语义建模。该方法结合联邦学习的分布式训练能力与先进的语言模型结构, 在保证隐私性的同时兼顾了推荐性能与通信效率。

在模型结构方面, 本文引入了具备解耦注意力机制与增强位置编码能力的 DeBERTaV3 作为本地客户端主干网络。与传统 BERT 模型相比, DeBERTaV3 能更精细地区分 Token 的语义信息与位置信息, 显著提升了长文本场景下的建模能力, 并在保持预训练权重共享的基础上实现个性化泛化。

在输入建模过程中, 采用 Token 表示机制将用户评论转化为结构化的语义单位, 有效保留了上下文信息与文本细粒度特征。Token 化不仅提升了模型对自然语言输入的处理能力, 也为后续的模型训练和跨客户端聚合提供了标准化输入形式。

在系统架构方面, 本文设计了一个中心协调式联邦学习框架, 模拟多个高校图书馆客户端之间的协同训练场景。在训练策略上, 采用冻结底层参数、仅微调 top-layer 与分类头的方式, 显著降低了通信开销, 提升了模型部署的可行性。训练过程中引入异步客户端参与机制, 以模拟真实环境下的设备异构性和在线不稳定性。

通过构建三组对照实验, 系统评估了不同模型结构与训练策略在推荐任务中的表现。实验结果表明, 本文方法在 Accuracy、F1、AUC 等指标上均取得优异表现, 尤其在 F1 值上达到 0.908, 几乎接近集中式训练的上限; 同时通信成本显著下降, 收敛轮次更快, 体现出极高的效率优势。

值得强调的是, 引入联邦学习后模型仅有约 1% 的性能下降, 但带来了显著的隐私保护收益与系统可扩展性。这表明联邦机制与强语义建模结构的结合, 是实现推荐系统隐私性与性能协同优化的有效路径。

综上所述, 本文提出的联邦推荐方法在保护用户隐私、适应异构数据环境和提升模型语义建模能力等方面均展现出优越性, 具有良好的理论价值与实际应用潜力。该方法为构建安全、高效的个性化推荐系统提供了可行方案。

基金项目

国家自然科学基金(62173231)。

参考文献

- [1] He, X., Liao, L., Zhang, H., Nie, L., Hu, X. and Chua, T. (2017) Neural Collaborative Filtering. *Proceedings of the 26th International Conference on World Wide Web*, Perth, 3-7 April 2017, 173-182. <https://doi.org/10.1145/3038912.3052569>
- [2] Zhang, Y. and Lai, G. (2017) DeepCoNN: Deep Cooperative Neural Networks for Sparse and High-Dimensional Recommender Systems. *Proceedings of the 10th ACM International Conference on Web Search and Data Mining*, Cambridge, 6-10 February 2017, 305-314.
- [3] Ammad-Ud-Din, M., Ivannikova, E., Khan, S.A., Oyomno, W., Fu, Q., Tan, K.E. and Flanagan, A. (2019) Federated Collaborative Filtering for Privacy-Preserving Personalized Recommendation System. <https://arxiv.org/abs/1901.09888>
- [4] He, P., Gao, J. and Chen, W. (2021) DeBERTav3: Improving DeBERTa Using ELECTRA-Style Pre-Training with Gradient-Disentangled Embedding Sharing. <https://doi.org/10.48550/arXiv.2111.09543>
- [5] McMahan, B., Moore, E., Ramage, D., Hampson, S. and Arcas, B.A. (2017) Communication-Efficient Learning of Deep Networks from Decentralized Data. <https://arxiv.org/abs/1602.05629>
- [6] 王超, 李斌, 刘金国. 一种改进的 Adam 优化算法及其在神经网络中的应用[J]. 计算机工程与应用, 2020, 56(24): 136-141.
- [7] 张昊, 刘挺. 深度学习中的损失函数综述[J]. 计算机科学, 2020, 47(6): 21-28.
- [8] 张海波, 魏秀参, 胡春明. 预训练语言模型的发展综述[J]. 软件学报, 2022, 33(4): 1132-1152.
- [9] Song, X., Salcianu, A., Song, Y., Dopson, D. and Zhou, D. (2021) Fast Wordpiece Tokenization. *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, Online and Punta Cana, November 2021, 2089-2103. <https://doi.org/10.18653/v1/2021.emnlp-main.160>
- [10] Sennrich, R., Haddow, B. and Birch, A. (2016) Neural Machine Translation of Rare Words with Subword Units. *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Berlin, August 2016, 1715-1725. <https://doi.org/10.18653/v1/p16-1162>
- [11] Hand, D.J. and Till, R.J. (2001) A Simple Generalisation of the Area under the ROC Curve for Multiple Class Classification Problems. *Machine Learning*, 45, 171-186. <https://doi.org/10.1023/a:1010920819831>