

王熙菲, 巫朝霞. 基于推荐系统的联邦协同变分自编码器[J]. 智能计算机与应用, 2025, 15(6): 73-80. DOI: 10.20169/j.issn.2095-2163.25040101

基于推荐系统的联邦协同变分自编码器

王熙菲, 巫朝霞

(新疆财经大学 统计与数据科学学院, 乌鲁木齐 830012)

摘要: 为了解决联邦学习中对用户数据隐私保护和数据稀疏问题, 提出了联邦协同变分自编码器模型(FCVAE)。该模型基于联邦学习框架, 各客户端采用变分自动编码器结构进行本地训练, 并引入迭代聚类方法增强模型表达能力。在 amazon_clothing, MovieLens-1M, MovieLens-100K 和 Amazon Office 四个数据集上的实验结果表明, FCVAE 在解释质量、个性化和一致性方面显著优于现有方法, 尤其在用户-项目交互稀疏的场景下, 展现出稳健的性能, 为推荐系统的进一步发展提供了有力支持。

关键词: 联邦学习; 变分自动编码器; 迭代聚类; 推荐系统; 隐私保护

中图分类号: TP301

文献标志码: A

文章编号: 2095-2163(2025)06-0073-08

Federated cooperative variational autoencoder based on recommender system

WANG Xifei, WU Zhaoxia

(School of Statistics and Data Science, Xinjiang University of Finance and Economics, Urumqi 830012, China)

Abstract: In order to solve the problem of user data privacy protection and data sparsity in federated learning, a Federated Cooperative Variational AutoEncoder model (FCVAE) is proposed. The model is based on the federated learning framework, and each client adopts the variational autoencoder structure for local training, and the iterative clustering method is introduced to enhance the expression ability of the model. Experimental results on four datasets, such as amazon_clothing, MovieLens-1M, MovieLens-100K, and Amazon Office, show that FCVAE is significantly better than the existing methods in terms of interpretation quality, personalization, and consistency, especially in the scenario of sparse user-project interaction, demonstrating robust performance. The further development of the recommender system provides beneficial support.

Key words: federated learning; variational autoencoders; iterative clustering; recommender systems; privacy protection

0 引言

在推荐系统中, 用户和项目之间的数据稀疏性一直是一个具有挑战性的研究课题。这种稀疏性导致了模型训练时的数据不足问题, 也限制了个性化推荐的准确性。传统的推荐算法, 如协同过滤和矩阵分解, 在处理高维稀疏数据时表现不佳, 很难有效地捕捉用户的兴趣模式。此外, 随着互联网应用的普及和个人隐私保护意识的提高, 如何在不暴露用户敏感信息的情况下提供高质量的推荐服务, 也是一个迫切需要解决的问题。目前, 联邦学习(Federated Learning, FL)^[1-2]得到了迅速发展, 作为

新兴的分布式机器学习框架, 能够有效应对并解决上述挑战。联邦学习允许多个参与方在不交换本地数据的情况下共同训练一个全局模型, 有效减少了数据外泄的可能性, 同时满足了法律监管的要求。

针对上述问题, 本文提出了联邦协同变分自编码器模型(Federated Learning Collaborative Variational AutoEncoder, FCVAE), 该模型不仅可以有效处理高维稀疏数据, 还通过引入 MoE 机制和迭代聚类技术来增强模型的表达能力, 更好地捕捉用户的兴趣分布。同时, 结合了联邦学习框架, 以降低通信成本并加强隐私保护。

作者简介: 王熙菲(1999—), 女, 硕士研究生, 主要研究方向: 联邦学习, 隐私保护。

通信作者: 巫朝霞(1975—), 女, 博士, 教授, 主要研究方向: 信息安全。Email: wuzhaoxia828@163.com。

收稿日期: 2025-04-01

哈尔滨工业大学主办 ◆ 学术研究与应用

1 相关工作

1.1 推荐系统研究现状

随着互联网的快速发展和用户可用数据量的不断增加,推荐系统(RS)^[3-4]技术已成为时下学术研究和开发的一个重要领域。推荐系统面临主要挑战之一是数据稀疏问题,特别是对于新用户或新项目而言。当可用数据不足时,传统的推荐算法往往难以提供准确的建议,从而导致用户体验不佳。尤其在新用户没有先前交互历史记录的场景中,这一问题将更加凸显,使得系统难以生成相关推荐^[5]。另一个重大挑战是冷启动问题,这是指向新用户推荐商品或是对缺乏用户评分的新商品来提供建议的困难。许多推荐系统严重依赖历史数据来进行预测,这一状况使得这个问题变得更加复杂。因此,新用户可能会收到与其偏好不符的通用推荐,从而导致困扰与脱离^[6]。将上下文信息集成到推荐系统中仍然是一项复杂的任务。用户的偏好可能因环境(例如时间、位置和社会影响)而有很大差异。然而,许多现有系统没有充分考虑这些背景因素,导致建议可能与特定情况不相关。应对这一挑战需要开发更复杂的模型,从而可以动态地将上下文信息纳入推荐过程^[7]。

1.2 联邦学习在VAE中的应用

联邦学习已成为机器学习的重要范例,特别是在医疗保健和生物医学应用领域。变分自动编码器(Variational AutoEncoder, VAE)^[8]是一类生成模型,已越来越多地与联邦学习集成,以提高其在各种应用中的性能,包括医学图像分析和患者数据建模,联邦学习FL与VAE的集成具有多种优势。

首先,FL与VAE的集成能够对去中心化数据进行VAE培训,这在高度重视数据隐私保护的医疗保健环境中至关重要。通过允许机构保留对其数据的控制权,同时仍为共享模型做出贡献,FL减轻了与数据共享相关的风险。研究表明,联邦VAE可以有效地从异构数据源学习表示,从而提高模型的鲁棒性和泛化性^[9]。大多数探索联邦VAE的研究都集中在医学成像任务上,特别是肿瘤学和放射学领域。例如,很大一部分研究集中于使用联邦VAE来执行肿瘤分割和分类等任务,其中从分布式数据集学习的能力可以显著提高模型性能。这些研究通常采用集中式聚合服务器来组合源自不同机构的模型参数,从而开发出受益于训练数据多样性的强大全局模型^[10-11]。

2 背景知识

2.1 联邦平均

联邦平均(Federated Averaging)是联邦学习中的一个重要概念。在联邦学习中,多个参与方各自训练本地模型,并将本地模型的参数更新发送到中央服务器,中央服务器根据这些参数调整来调整全局模型。以下是联邦平均算法:

设共有 K 个客户端,中心服务器初始化模型参数,执行若干轮(round),每轮选取至少1个、至多 K 个客户端参与训练,接下来每个被选中的客户端同时在自己的本地根据服务器下发的本轮(t 轮)模型 w_t 用自己的数据训练自己的模型 w_{t+1}^k ,上传回服务器。服务器将收集来的各客户端的模型根据各方样本数量用加权平均的方式进行聚合,得到下一轮的模型 w_{t+1} :

$$w_{t+1} \leftarrow \sum_{k=1}^K \frac{n_k}{n} w_{t+1}^k \quad (1)$$

其中, n_k 表示客户端 k 上的样本数量, n 表示所有被选中客户端的总样本数量。

2.2 变分自编码器

VAE^[12-13]模型是一种生成模型,通过神经网络训练得到推断网络和生成网络,含有隐变量,可以生成不包含在输入数据中的新数据。这一模型在生成多种复杂数据方面显示出了巨大潜力,包括手写数字图像、人脸图像、门牌号图像、CIFAR 图像、物理场景模型、分割图像以及从静态图像进行预测等。VAE的核心贡献在于将生成模型的学习与变分推断相结合,提出一种高效的训练方法,以解决传统生成模型中的计算难题。首先,通过引入变分分布来近似后验分布,VAE使得生成模型的训练能够通过标准的梯度下降方法进行优化,而无需复杂的采样或优化算法。其次,为了解决梯度反向传播中的不可微问题,VAE采用重参数化技巧,使潜在变量的采样过程变得可微分,从而实现高效的训练。最后,通过自编码器的结构,VAE能够端到端地进行训练,将数据的编码与解码过程结合起来,同时进行潜在空间的学习和生成任务的优化。这些关键技术的结合使得VAE在生成模型的研究中展现出了很大的潜力,为机器学习领域带来了新的发展方向。

2.3 混合专家

混合专家系统(MoE)^[14]是一种神经网络模型,将不同数据产生方式的数据集进行有效整合。与传统神经网络不同的是,MoE通过分离训练多个专家

模型来处理数据,然后通过门控模块选择合适的专家模型进行输出。这些专家模型可以采用不同的函数,包括各种线性和非线性函数,最终通过加权组合得到整体输出。MoE 的优势在于将多个模型整合到一个任务中,提高了模型的表现力和适应性。

MoE 架构通过引入专家网络层和门控机制,实现了对输入空间的有效分割和路由选择,进而提高了模型的效率和性能。稀疏模型和组合设计让

MoE 架构在处理非线性监督学习问题时具有明显的优势,为解决复杂任务提供了强大的工具和方法。

3 联邦协同变分自编码器模型方案设计

本方案的核心是利用迭代聚类机制对个性化项目进行分组。为后续对算法描述方便的需要,现将后续章节将要使用到的参数汇总见表 1。

表 1 符号设置
Table 1 Symbol settings

符号	参数
$D = \{x^{(u)}\}_{u=1}^N$	训练集,包含 N 个用户的用户-物品交互向量
$x^{(u)}$	用户 u 的用户-物品交互向量
d_z	潜在空间维度
likelihood	输出层使用的似然函数类型
K	聚类数,表示将用户分为多少类别进行兴趣聚合
R	聚类迭代次数
E	专家数量,在解码器部分使用的混合专家模块中的专家数目
β	正则化参数
η	学习率
T	训练轮次
$\mu^{(u)}, \log \sigma^{(u)}$	对于用户 u ,由编码器生成的潜在变量的均值和对数方差
$z^{(u)}$	用户 u 的潜在变量
m_k	第 k 类别的原型向量,代表该类别的中心点
$A^{(u)}$	注意力矩阵,描述了用户 u 对不同类别的注意力权重
$\hat{x}^{(u)}$	用户 u 的重建用户-物品交互向量
ϵ	标准正态分布中的随机噪声,用于实现重参数化技巧
L_{ll}	似然损失,衡量模型预测值与真实值之间的差异
L_{kl}	KL 散度损失,衡量潜在变量分布与标准正态分布之间的距离
$\nabla_{\theta} L$	损失函数关于模型参数 θ 的梯度,用于指导模型参数的更新方向

3.1 算法设计

3.1.1 迭代聚类

迭代聚类^[15]是模型中的关键部分,是通过动态更新用户兴趣表示来增强模型对用户行为的理解。这里,对各流程步骤给出阐释分述如下。

(1)初始化原型。首先,为每个类别(或簇)初始化一个原型向量 $m_k \in \mathbb{R}^{d_{input}}$,其中 $k = 1, \dots, K$ 表示类别索引, d_{input} 表示输入数据的维度。这些原型

向量代表了不同类别中心的初始估计。

(2)计算注意力矩阵。对于每一个用户 u 的交互向量 $x^{(u)}$,计算其与所有原型向量之间的相似度,并将其归一化以形成注意力矩阵 $A^{(u)}$ 。

归一化输入 $x^{(u)}$ 和原型的目的是为了确保计算相似度时不会因为尺差异导致偏差。具体来说,对于用户 u 的交互向量 $x^{(u)}$ 和第 k 类别的原型 m_k ,首先分别计算对应的 $L2$ 范数 $\|x^{(u)}\|_2$ 和 $\|m_k\|_2$ 。

然后,基于内积和归一化结果,使用 Gumbel-Softmax 函数计算注意力矩阵 $\mathbf{A}^{(u)}$ 中的元素 $A_{ik}^{(u)}$:

$$A_{ik}^{(u)} = \text{Gumbel-Softmax}\left(\frac{\langle \mathbf{x}_i^{(u)}, \mathbf{m}_k \rangle}{\tau \|\mathbf{x}_i^{(u)}\| \cdot \|\mathbf{m}_k\|}\right) \quad (2)$$

其中, τ 表示温度参数,用于控制分布的平滑程度。较低的 τ 值会使分配更接近于 one-hot 形式,而较高的值会产生更加分散的概率分布。

(3)更新原型。在获得注意力矩阵后,根据用户的交互向量及其对应的注意力权重更新每个类别的原型向量 $\mathbf{m}_k$ 。这个过程可以通过下式来实现:

$$\mathbf{m}_k \leftarrow \frac{\sum_u A_{uk} \mathbf{x}^{(u)}}{\sum_u A_{uk}} \quad (3)$$

(4)聚合用户兴趣。根据更新后的原型向量,再次计算注意力矩阵,并据此聚合用户的兴趣表示。具体地,对于每个用户 u ,其新的兴趣表示 $\mathbf{x}_k^{(u)}$ 可以通过下式来计算:

$$\mathbf{x}_k^{(u)} = \sum_{k=1}^K A_{uk} \mathbf{x}^{(u)} \quad (4)$$

其中, $\mathbf{x}_k^{(u)}$ 表示用户 u 的交互向量经过加权求和后的新表示,权重 A_{uk} 来自最新的注意力矩阵。

(5)重复步骤(2)~(4)。

上述过程进行 R 轮迭代,每一轮迭代都会更新原型向量,并根据更新后的原型重新计算注意力矩阵,从而逐步细化用户的兴趣表示和类别原型。

3.1.2 编码

编码部分是将用户兴趣表示从输入空间映射到潜在空间。这里对各步骤将展开探讨论述如下。

(1)输入预处理。在开始编码之前,首先需要对用户的交互向量进行适当的预处理。这通常涉及到通过一个线性变换层来减少输入维度。假设原始用户-物品交互矩阵为 $\mathbf{x}^{(u)}$,则经过降维操作后的表示为:

$$\mathbf{x}_k^{(u)} = \text{reduce_dim}(\mathbf{x}^{(u)}) \quad (5)$$

其中, reduce_dim 表示一个线性变换,用于将输入数据映射到一个较低维度的空间,以便于后续处理。

(2)编码器网络。编码器由一系列全连接层组成,每一层包括线性变换和激活函数。具体来说,对于每一层 l ,其输出可以表示为:

$$\mathbf{h}^{(l)} = \text{act}(\mathbf{W}^{(l)} \mathbf{h}^{(l-1)} + \mathbf{b}^{(l)}) \quad (6)$$

其中, $\mathbf{h}^{(l)}$ 表示第 l 层的隐藏表示; act 表示激活函数; $\mathbf{W}^{(l)}$ 和 $\mathbf{b}^{(l)}$ 分别表示该层的权重矩阵和偏

置向量。

编码器的最后一层并不包含激活函数,而是直接输出 2 个向量、即:均值向量 $\boldsymbol{\mu}^{(u)}$ 和对数方差向量 $\log \boldsymbol{\sigma}^{(u)}$ 。这 2 个向量分别定义了潜在变量分布的均值和方差,即:

$$\boldsymbol{\mu}^{(u)}, \log \boldsymbol{\sigma}^{(u)} = \text{Encoder}(\mathbf{x}_k^{(u)}) \quad (7)$$

(3)重参数化技巧。为了使得潜在变量 $\mathbf{z}^{(u)}$ 可以通过梯度下降方法优化,使用了重参数化技巧。根据均值向量 $\boldsymbol{\mu}^{(u)}$ 和对数方差向量 $\log \boldsymbol{\sigma}^{(u)}$,通过下式采样得到潜在变量:

$$\mathbf{z}^{(u)} = \boldsymbol{\mu}^{(u)} + \boldsymbol{\epsilon} \cdot \exp(0.5 \cdot \log \boldsymbol{\sigma}^{(u)}), \boldsymbol{\epsilon} \sim N(0,1) \quad (8)$$

其中, $\boldsymbol{\epsilon}$ 表示从标准正态分布中采样的随机噪声。这个过程允许计算损失相对于模型参数的梯度,从而实现端到端的训练。

(4)输出。最终,编码器输出的是每个用户的潜在变量表示 $\mathbf{z}^{(u)}$ 。

3.1.3 使用专家混合解码

使用专家混合解码增强了模型的表达能力,使得解码器可以根据不同的输入动态选择不同的“专家”进行处理。这里,对各步骤进行具体阐述如下。

(1)解码器网络。首先,从潜在空间中的表示开始,通过一系列全连接层将其映射回原始输入空间的维度。这一过程中,每一层包括线性变换和激活函数。具体来说,对于每一层 l ,其输出可以表示为:

$$\mathbf{h}^{(l)} = \text{act}(\mathbf{W}^{(l)} \mathbf{h}^{(l-1)} + \mathbf{b}^{(l)}) \quad (9)$$

解码器的最后一层不包含激活函数,直接输出一个中间表示:

$$\mathbf{h}^{(u)} = \text{Decoder}(\mathbf{z}^{(u)}) \quad (10)$$

(2)计算专家输出。在 MoE 机制中,存在多个专家 (E 个),每个专家都是一个独立的神经网络层,负责生成候选的重建结果。对于每个专家 e ,基于中间表示 $\mathbf{h}^{(u)}$ 计算其输出 $\mathbf{h}_e^{(u)}$:

$$\mathbf{h}_e^{(u)} = \text{Expert}_e(\mathbf{h}^{(u)}) \quad (11)$$

其中, Expert_e 表示第 e 个专家的计算过程,通常也是一个线性变换加上可能的激活函数。

(3)计算门控值。为了决定每个专家对最终输出的贡献程度,需要引入一个门控网络来计算各个专家的权重。门控网络接收中间表示 $\mathbf{h}^{(u)}$ 作为输入,并输出一个概率分布,表示每个专家被选中的概率。具体地,通过以下公式计算门控值 $\mathbf{g}_e^{(u)}$:

$$\mathbf{g}_e^{(u)} = \text{Softmax}(\text{Gate}(\mathbf{h}^{(u)})) \quad (12)$$

其中, Gate 表示门控网络,通常由一个或多个

全连接层组成; *Softmax* 函数用于将门控网络的输出转换为一个有效的概率分布, 确保所有专家的概率之和为 1。

(4) 组合专家输出。根据门控值, 将各个专家的输出加权求和, 得到最终的重建用户-物品交互向量 $\hat{\mathbf{x}}^{(u)}$ 。具体而言, 通过下式实现:

$$\hat{\mathbf{x}}^{(u)} = \sum_{e=1}^E \mathbf{g}_e^{(u)}, \mathbf{h}_e^{(u)} \quad (13)$$

(5) 输出。最终, 经过上述过程后得到的 $\hat{\mathbf{x}}^{(u)}$ 就是用户 u 的重建用户-物品交互向量。这个向量应该尽可能接近原始输入, 以便于后续计算重构误差并优化模型参数。

3.1.4 损失计算

损失计算决定了模型参数如何更新以最小化重建误差并正则化潜在变量分布, KL 散度损失确保潜在变量 $z(u)$ 的分布接近标准正态分布, 从而增加模型的泛化能力。

KL 散度损失衡量了编码器输出的潜在变量分布 $q(z | x)$ 和先验分布 $p(z)$ 之间的差异。具体公式如下:

$$L_{\text{kld}} = -\frac{1}{2} \sum_j [1 + \log(\sigma_j^{(u)}) - (\mu_j^{(u)})^2 - (\sigma_j^{(u)})^2] \quad (14)$$

其中, $\mu^{(u)}$ 和 $\sigma^{(u)}$ 分别表示编码器输出的均值和对数方差。

3.1.5 损失优化

(1) 反向传播与梯度下降: 计算损失相对于模型参数的梯度, 并使用优化算法 (如 Adam) 更新模型参数。

(2) 循环执行: 损失计算和参数更新步骤在整个训练集上重复进行多次 (由 T 决定), 直到模型收敛或达到预设的停止条件。

总而言之, 多门控机制不仅可以捕捉到用户的复杂行为模式, 还能为不同类型的用户提供高度个性化的解释服务。这种个性化定制不仅可以提高用户满意度, 还能优化用户体验, 使用户更容易接受和理解所提供的服务。伪代码见算法 1。

算法 1 FCVAE 算法

输入 训练数据集 D 、 Z_{dim} 、*likelihood*、 K 、 τ 、 β 、 E 、 η 、 T

输出 训练好的模型参数 θ 、 m_k

1. 初始化: *ae_structure*、*act_fn*

2. For each epoch $t = 1, \dots, T$:

3. Shuffle the training data and split into mini-batches of size B

4. For each mini-batch of user data $\{\mathbf{x}^{(u)}\}_{u=1}^B$

5. 迭代聚类

6. $\mathbf{x}_k^{(u)} = \sum_{k=1}^K A_{uk} \mathbf{x}^{(u)}$

7. 编码

8. $\mathbf{z}^{(u)} = \mu^{(u)} + \epsilon \cdot \exp(0.5 \cdot \log \sigma^{(u)})$, $\epsilon \sim N(0, 1)$

9. 解码

10. $\hat{\mathbf{x}}^{(u)} = \sum_{e=1}^E \mathbf{g}_e^{(u)}, \mathbf{h}_e^{(u)}$

11. 损失计算

3.2 框架设计

该模型结合了 MoE 机制和迭代聚类技术, 以增强对用户兴趣分布的学习能力。具体来说, 变分自编码器通过编码器将原始输入映射到低维潜在空间, 并利用 MoE 架构在解码阶段实现灵活输出; 同时, 迭代聚类算法帮助模型动态学习用户的多模态兴趣簇, 提高了个性化推荐的准确性。此外, 将此模型与联邦学习框架相结合, 以应对数据安全性和通信成本的挑战。

以下将详细介绍 FCVAE 的模型部署结构及其工作原理, 工作流程如图 1 所示。在 FCVAE 架构中, 模型部署结构主要分为中央服务器和客户端两部分。中央服务器负责模型的初始化、分发、收集和聚合。客户端代表了不同终端设备或服务提供者参与联邦学习过程, 同时也持有各自的私有数据集但不直接共享数据。客户端的主要任务包括本地训练和共享权重更新。在本地训练过程中, 客户端利用本地数据对接收到的全局模型进行训练, 采用迭代聚类技术和 MoE 结构增强模型对复杂多模态数据的理解能力。同时, 通过 FCVAE 模型中的迭代聚类技术优化用户兴趣表示和推荐的相关性。完成一定数量的本地训练后, 客户端将计算出的模型参数变化发送回中央服务器, 确保用户隐私得到有效保护的同时促进全局模型的不断进化。

通过上述设计, FCVAE 不仅有效地解决了联邦学习推荐系统中用户-项目交互数据稀疏的问题, 同时提高了推荐系统的准确性和鲁棒性。

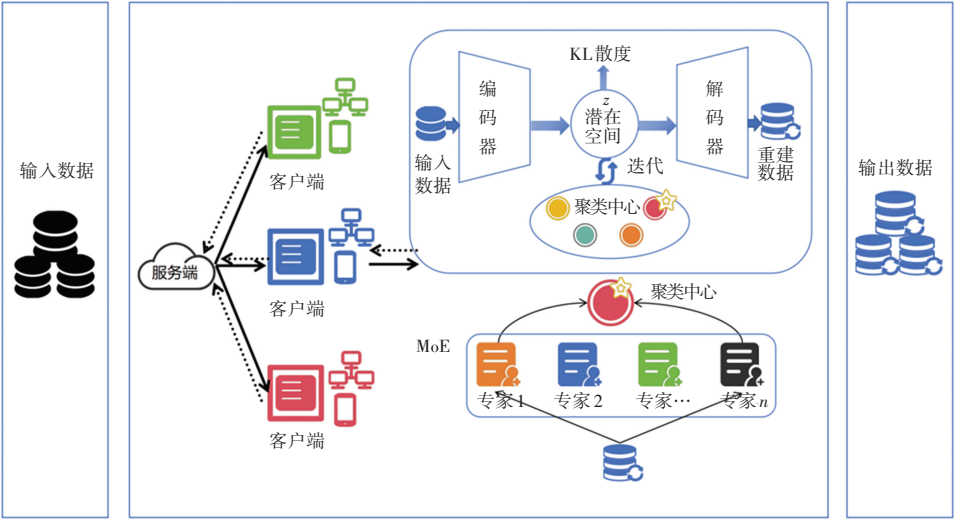


图 1 FCVAE 模型结构

Fig. 1 Structure of the FCVAE model

4 实验

本文实验中模型训练所使用的硬件配置具体如下:NVIDIA GeForce RTX 4070 显卡和 AMD Ryzen 7 5800H 处理器,内存为 32 G,所有实验均采用 Pytorch 框架完成。

4.1 实验设置

4.1.1 数据集

在本次研究中,选用了 4 个公开可用的数据集进行实验: Amazon _ Clothing、MovieLens - 1M、MovieLens-100K 和 Amazon Office Products。细节详见如下。

(1) Amazon Clothing :该数据集包含了丰富的用户评论和评分信息,涵盖了大量亚马逊服装产品。

(2) MovieLens - 1M 和 MovieLens - 100K: 是 2 个电影评分数据集,分别收录了约一百万和 10 万条用户评分。这些数据集为研究人员提供了丰富的信息,帮助其准确了解用户对电影的评价和喜好。

(3) Amazon Office Products:是一个专注于办公用品类别的数据集。

对于每个数据集,采用相同的策略进行了划分,将 10%的数据作为测试集,然后从剩余 90%的数据中再次抽取 10%作为验证集。此外,还定义了一个二元偏好阈值,即只有当用户对某项办公用品的评分高于此值时,才会认为用户对该项产品感兴趣。

4.1.2 参数设置

针对该模型的具体实现,设置了以下超参数来优化其性能,见表 2。

表 2 实验参数设置

Table 2 Experimental parameter settings

参数名称	设置
隐含层维度 (k)	256
神经元数量	20
激活函数	tanh
似然函数类型	mult
轮次 ($epoch$)	50
$batch_size$	128
学习率	0.001
KL 散度	0.01

4.1.3 基线模型

为了深入评估 FCVAE 在推荐系统中的有效性,实验将与如下基线模型进行了综合性能比较:

(1)VAECF。是一种通过变分自编码器学习用户-项目交互矩阵潜在表示的算法,通过这种方式生成推荐。

(2)Cvae^[16]。在 VAECF 基础上引入了条件信息,如用户属性或项目属性,以进一步提高推荐结果的准确性和相关性。

(3)Revae^[17]:通过引入正则化项来减少变分自编码器的过拟合问题,从而提高模型的泛化能力。

(4)Bivaec^[18]。一种同时考虑用户和项目潜在表示的双向变分自编码器,用以改进推荐效果。

(5)BVAE^[19]。提出了一个行为感知的变分自编码器(BVAE),专为多任务推荐设计,能够同时处理用户与项目之间的多种行为关系。

(6)MRVAE^[20]。扩展了传统的协同过滤方法,引入多样关系建模能力,能够处理用户与项目之间不同类型和复杂度的交互关系。

(7)VMoSE^[21]。结合了变分混合专家和多模态学习,提出了一个能够从不同数据源(如文本、图像、行为日志)中提取信息的自表达式编码器。

4.1.4 评价指标

为了衡量不同模型的表现,采用了 4 个推荐系统评估指标: *MRR* (Mean Reciprocal Rank) 用于确定推荐列表中正确项目的排名位置,数值越高表示排序质量越好。*NDCG* 反映了前 20 个推荐项的相关性得分,考虑了位置权重。*Recall* 衡量在前 20 个推荐项中用户真正感兴趣的项目所占比例。*Precision* 表明前 20 个推荐项中有多少是真正被用户喜欢的。为了验证 FCVAE 模型的性能,在数据集上进行了多次的对比实验。实验结果见表 3~表 6。

表 3 FCVAE 模型在 MovieLens-1M 数据集对比结果
Table 3 Comparison results of FCVAE models in the MovieLens-1M dataset

数据集	算法	<i>MRR</i>	<i>NDCG</i>	<i>Precision</i>	<i>Recall</i>
MovieLens-1M	VAECF	0.398 7	0.207 5	0.129 9	0.189 7
	Revae ^[17]	0.398 9	0.210 5	0.130 2	0.180 5
	Cvae ^[16]	0.412 5	0.220 3	0.151 7	0.220 4
	Bivaecf ^[18]	0.420 5	0.240 6	0.150 7	0.237 8
	VMoSE ^[21]	0.418 3	0.224 1	0.1403	0.209 6
	MRVAE ^[20]	0.343 9	0.195 9	0.122 2	0.210 1
	BVAE ^[19]	0.264 9	0.120 2	0.077 6	0.111 3
	FCVAE	0.452 5	0.243 0	0.156 5	0.207 6

表 4 FCVAE 模型在 MovieLens-100K 数据集对比结果
Table 4 Comparison results of FCVAE models in the MovieLens-100K dataset

数据集	算法	<i>MRR</i>	<i>NDCG</i>	<i>Precision</i>	<i>Recall</i>
MovieLens-100K	VAECF	0.228 1	0.138 8	0.065 1	0.201 2
	Revae ^[17]	0.250 6	0.139 2	0.060 5	0.201 5
	Cvae ^[16]	0.251 2	0.145 9	0.067 1	0.220 8
	Bivaecf ^[18]	0.261 9	0.160 7	0.068 2	0.249 8
	VMoSE ^[21]	0.205 5	0.138 6	0.064 0	0.207 1
	MRVAE ^[20]	0.206 0	0.132 7	0.060 1	0.204 7
	BVAE ^[19]	0.168 8	0.103 5	0.044 5	0.160 2
	FCVAE	0.277 9	0.181 2	0.080 8	0.257 0

表 5 FCVAE 模型在 Amazon_Office 数据集对比结果
Table 5 Comparison results of FCVAE models in the Amazon_Office datasets

数据集	算法	<i>MRR</i>	<i>NDCG</i>	<i>Precision</i>	<i>Recall</i>
Amazon_Office	VAECF	0.043 4	0.036 0	0.008 2	0.071 2
	Revae ^[17]	0.042 3	0.035 8	0.009 1	0.071 6
	Cvae ^[16]	0.043 5	0.041 2	0.009 5	0.075 9
	Bivaecf ^[18]	0.048 9	0.042 5	0.009 7	0.081 2
	VMoSE ^[21]	0.048 8	0.041 3	0.009 3	0.077 5
	MRVAE ^[20]	0.045 0	0.035 5	0.008 0	0.066 7
	BVAE ^[19]	0.014 1	0.009 9	0.002 1	0.018 4
	FCVAE	0.050 5	0.042 7	0.009 5	0.081 7

表 6 FCVAE 模型在 Amazon_Clothing 数据集对比结果
Table 6 Comparison results of FCVAE models in the Amazon_Clothing datasets

数据集	算法	<i>MRR</i>	<i>NDCG</i>	<i>Precision</i>	<i>Recall</i>
Amazon_Clothing	VAECF	0.028 8	0.037 0	0.004 5	0.085 3
	Revae ^[17]	0.030 5	0.038 0	0.004 6	0.086 1
	Cvae ^[16]	0.034 5	0.040 8	0.005 1	0.086 4
	Bivaecf ^[18]	0.054 0	0.045 9	0.005 3	0.089 8
	VMoSE ^[21]	0.026 7	0.034 8	0.004 4	0.080 3
	MRVAE ^[20]	0.030 6	0.036 4	0.003 6	0.067 6
	BVAE ^[19]	0.026 6	0.033 0	0.004 1	0.073 3
	FCVAE	0.064 0	0.074 3	0.006 9	0.127 0

4.2 消融实验

为了更深刻地理解 FCVAE 模型的各个组成部分的作用,将 FCVAE 与 2 个变体模型进行对比: FL+MoE(仅保留混合专家模型)和 FL+MultiGroup(仅保留迭代聚类结构)。通过这种方式,可以评估每个组件对整体性能的影响,并进一步确认 FCVAE 的优势所在。

在本节的消融实验中,继续使用前文提到的 4 个数据集 (Amazon Clothing、MovieLens - 1M、MovieLens-100K 和 Amazon Office Products),并采用相同的评估指标 (*MRR*、*NDCG*、*Recall* 和 *Precision*) 来衡量模型的表现。所有实验均在同一环境下运行,以确保结果的可比性。

消融实验结果对比见表 7。根据表 7 数据可以得知,FCVAE 在 4 个数据集上的表现都优于其变体模型 FL+MoE 和 FL+MultiGroup,这意味着 FCVAE 的每个组件都对性能有积极的贡献。综上所述,消融实验结果进一步验证了 FCVAE 框架的有效性和先进性。每个组件都在提升模型性能方面发挥了重

要作用,而这些机制的协同作用则使得 FCVAE 成为处理非独立同分布数据和通信成本问题的理想选择。

表 7 消融实验结果对比

Table 7 Comparison of ablation experimental results					
数据集	算法	MRR	NDCG	Precision	Recall
MovieLens-1M	FL+MoE	0.434 0	0.235 7	0.146 2	0.220 9
	FL+MultiGroup	0.441 0	0.241 0	0.152 3	0.201 1
	FCVAE	0.452 5	0.243 0	0.156 5	0.207 6
MovieLens-100K	FL+MoE	0.227 7	0.150 0	0.068 2	0.226 3
	FL+MultiGroup	0.234 7	0.171 6	0.074 6	0.266 0
	FCVAE	0.277 9	0.181 2	0.080 8	0.257 0
Amazon_Office	FL+MoE	0.044 8	0.037 3	0.008 6	0.073 6
	FL+MultiGroup	0.049 5	0.041 2	0.009 2	0.080 7
	FCVAE	0.050 5	0.042 7	0.009 5	0.081 7
Amazon_Clothing	FL+MoE	0.057 2	0.071 1	0.007 7	0.142 0
	FL+MultiGroup	0.043 6	0.054 1	0.005 5	0.101 0
	FCVAE	0.064 0	0.074 3	0.006 9	0.127 0

5 结束语

本文提出了一种创新联邦学习框架 FCVAE,解决传统方法的通信成本和模型性能问题,特别适用于推荐系统。通过结合变分自编码器、混合专家模型和迭代聚类技术,显著提升了准确性和个性化程度,并在稀疏数据处理方面表现优异。框架支持多客户端协同训练,无需共享本地数据,保护隐私并降低数据泄露风险。此外,通过迭代聚类+MoE 结构,捕获复杂多模态数据特征,提高了模型灵活性和表达能力。实验结果显示,在 4 个真实数据集上,FCVAE 优于基线方法,具备广泛应用潜力,可扩展至医疗健康、金融风控等领域,并通过优化模型机制降低通信成本,支持大规模部署。

参考文献

[1] SINGH P, SINGH M, SINGH R, et al. Federated Learning: Challenges, methods, and future directions [J]. Federated Learning for IoT Applications, 2022, 37(3): 50–60.

[2] 肖雄,唐卓,肖斌,等. 联邦学习的隐私保护与安全防御研究综述[J]. 计算机学报, 2023, 46(5):1019–1044.

[3] BATMAZ Z, YUREKLI A, BILGE A, et al. A review on deep learning for recommender systems: challenges andremedies[J]. Artificial Intelligence Review, 2019, 52(1): 1–37.

[4] CHEN Junyang, ZOU Guoxuan, ZHOU Pan, et al. Sparse enhanced network: An adversarial generation method for robust augmentation in sequential recommendation[C]//Proceedings of

the Thirty – Eighth AAAI Conference on Artificial Intelligence. Vancouver, Canada;AAAI, 2024: 8283–8291.

[5] WAYESA F, LERANSO M, ASEFA G, et al. Pattern – based hybrid book recommendation system using semantic relationships [J]. Scientific Reports, 2023,13(1): 32–44.

[6] BALCHANOWSKI M, BORYCZKA U. A comparative study of rank aggregation methods in recommendation systems [J]. Entropy, 2023,25(1): 132–151.

[7] FENG Yilin. Enhancing e – commerce recommendation systems through approach of buyer’s self – construal: necessity, theoretical ground, synthesis of a six – step model, and research agenda[J]. Frontiers in Artificial Intelligence, 2023,6: 1167735.

[8] SINGH A, OGUNFUNMI T. An overview of variational autoencoders for source separation, finance, and bio – signal applications[J]. Entropy, 2021,24(1): 55–110.

[9] SALMERON J L, AREVALO I, RUIZ – CELMA A. Benchmarking federated strategies in Peer – to – Peer Federated learning for biomedicaldata[J]. Heliyon, 2023,9(6): e16925.

[10]CROWSON M G, MOUKHEIBER D, AREVALO A R, et al. A systematic review of federated learning applications for biomedical data[J]. PLOS Digit Health, 2022,1(5): e0000033.

[11] ELTAGER M, ABDELAAL T, CHARROUT M, et al. Benchmarkingvariational AutoEncoders on cancer transcriptomics data[J]. PLoS One, 2023,18(10): e0292126.

[12]KINGMA D P, WELLING M. Auto – encoding variational bayes [J]. arXiv preprint arXiv,1312.6114v1,2013.

[13] DOERSCH C. Tutorial onvariational autoencoders [J]. arXiv preprint arXiv,1606.05908,2016.

[14] LEPIKHIN D, LEE H, XU Y, et al. GShard: Scaling giant models with conditional computation and automatic sharding[J]. arXiv preprint arXiv, 2006.16668,2020.

[15]TANG F, SHEN Y, ZHANG H, et al. GaVaMoE: Gaussian – variational gated mixture of experts for explainable recommendation[J]. arXiv preprint arXiv, 2410.11841,2024.

[16] KIM J, KONG J, SON J. Conditional variational autoencoder with adversarial learning for end – to – end text – to – speech [J]. arXiv preprint arXiv, 2106.06103,2021.

[17] JOY T, SCHMON S, TORR P, et al. Rethinking semi – supervised learning in VAEs [J]. arXiv preprint arXiv, 2006.10102,2020.

[18]TRUONG Q T, SALAH A, LAUW H W. Bilateral variational autoencoder for collaborative filtering [C]//Proceedings of the 14th ACM International Conference on Web Search and Data Mining. New York;ACM, 2021: 292–300.

[19] RAO Q, PAN W, ZHONG M. BVAE: Behavior – aware variational autoencoder for multi – behavior multi – task recommendation[C]//Proceedings of the 17th ACM Conference on Recommender Systems. New York;ACM, 2023: 625–636.

[20] PAN Zhou, LIU Wei, YIN Jian. MRVAE: Variational autoencoder with multiple relationships for collaborative filtering [C]//Proceedings of the Web Engineering. Cham; Springer, 2022: 16–30.

[21] YI Jing, CHEN Zhenzhong. Variational mixture of stochastic experts auto – encoder for multi – modal recommendation[J]. IEEE Transactions on Multimedia, 2024, 26: 8941–8954.