

可解释推荐的混合方法：结合合成坐标、Kolmogorov-Arnold Network和Transformer

艾 均, 徐戈辉, 苏 湛, 苏 杭

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4月1日; 录用日期: 2025年5月1日; 发布日期: 2025年5月8日

摘 要

近年来, 可解释推荐因其不仅能提供个性化推荐, 还能解释推荐原因而备受关注。基于Transformer的文本生成技术尽管可以生成高度自然的推荐解释, 但其高模型复杂性和计算成本也带来了挑战。为应对这些问题, 本文提出了一种基于Transformer架构的混合可解释推荐模型。该模型结合了合成坐标和Kolmogorov-Arnold网络(KAN), 通过将用户和物品映射至低维嵌入空间, 利用空间关系引导解释生成, 实现了高效的评分预测和解释生成。实验结果表明, 该模型在评分预测和解释生成任务中均达到了最先进的水平, 且在BLEU-4和USR等关键指标上取得了显著提升。研究还揭示了通过优化嵌入向量和模块化设计的结合, 能够有效提升可解释推荐系统的性能, 同时降低模型复杂性、计算难度及资源消耗。

关键词

可解释推荐, 合成坐标, Kolmogorov-Arnold Network, Transformer

A Hybrid Approach to Explainable Recommendation: Combining Synthetic Coordinate, Kolmogorov-Arnold Network and Transformer

Jun Ai, Gehui Xu, Zhan Su, Hang Su

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

Received: Apr. 1st, 2025; accepted: May 1st, 2025; published: May 8th, 2025

Abstract

Explainable recommendations have gained attention for offering both personalized suggestions

文章引用: 艾均, 徐戈辉, 苏湛, 苏杭. 可解释推荐的混合方法: 结合合成坐标、Kolmogorov-Arnold Network 和 Transformer [J]. 建模与仿真, 2025, 14(5): 13-30. DOI: 10.12677/mos.2025.145370

and clear explanations. While Transformer-based text generation produces natural explanations, it introduces challenges of model complexity and computational cost. This paper proposed a hybrid model combining synthetic coordinates and Kolmogorov-Arnold Networks (KAN) with Transformer architecture to address these issues. By mapping users and items into a low-dimensional space and leveraging spatial relationships, the model achieved efficient rating prediction and explanation generation. Experimental results showed state-of-the-art performance, with significant improvements in BLEU-4 and USR metrics. The study demonstrates that optimizing embeddings and using a modular design enhance system performance while reducing complexity and computational difficulty and resource consumption.

Keywords

Explainable Recommendation, Synthetic Coordinates, Kolmogorov-Arnold Network, Transformer

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

推荐系统(Recommender System) [1]是一种通过分析用户行为数据、偏好和历史记录,向用户推荐其可能感兴趣的内容或商品的技术。其作用是帮助用户快速发现他们感兴趣的信息或商品,提高用户体验,并增加用户与平台的互动,并且在推动商业增长和提升客户满意度方面也具有重要意义。

可解释推荐(Explainable Recommendation) [2]-[4]是推荐系统领域中的一个重要分支,旨在不仅提供个性化推荐,还能够向用户解释推荐背后的原因。对于用户而言,可解释推荐能帮助用户更好地理解为何会得到某些推荐,从而提高用户满意度,同时也使平台能够通过透明化的推荐过程来优化算法和提升服务质量。因此,近年来可解释推荐受到研究人员和工业界的广泛关注。

在可解释推荐任务中,引导式文本生成(Guided Text Generation) [5]-[7]是一种常用的技术。该技术通过特定的提示或相关信息,引导模型的生成方向,从而生成符合特定需求的文本。该技术能够生成高质量且具有针对性的内容,确保输出内容更符合用户需求。此外,该技术减少了生成过程中的不确定性,使生成过程更加可控和一致,从而在多种应用场景中实现更高的灵活性和生成质量。

近年来,在可解释推荐有如下三点主要的发展趋势:

第一,使用自然语言生成技术,生成模型如 Transformer、变分自编码器(VAE)和 Diffusion 模型在解释生成中发挥了重要作用。Transformer 模型以其强大的序列建模能力生成流畅的解释[8]-[10]; VAE 通过潜在空间建模支持多样化的解释[11] [12]; Diffusion 模型则通过逐步生成的过程提升解释质量[13] [14]。这些技术的结合使得推荐系统能够生成更人性化和精准的解释。

第二,大语言模型(LLMs) [15]如 GPT [16]-[18]和 BERT [19]在可解释推荐中得到越来越广泛的应用。这些模型通过大规模文本数据中的语言模式和语义关系,生成自然且连贯的解释[20]-[27]。特别是在复杂的推荐任务中,LLMs 能够结合上下文信息,为用户提供更加精准和深入的解释,帮助用户更好地理解推荐的依据和原因。

第三,为了提高解释的全面性和准确性,研究者们逐渐将知识图谱、方面分析和多模态信息结合起来。知识图谱提供了丰富的实体和关系信息,帮助解释融入更多背景知识和语境[28]-[30]; 方面分析聚焦于用户的具体需求和偏好,生成更贴合用户需求的解释[31]-[34]; 多模态信息整合了文本、图像和声音等

数据形式, 使推荐解释更加丰富且个性化[35]。

尽管可解释推荐的相关研究已经取得了一定的进展, 但目前仍面临着诸多挑战: 1) 计算成本与模型效果的权衡: 高效模型可能牺牲效果, 复杂模型虽效果佳但计算成本高[2]; 2) 模型复杂性与可解释性的矛盾: 深度模型虽精度高, 但难以解释, 而简化模型则可能影响性能[3]; 3) 衡量指标: 传统指标如 BLEU 难以全面评估模型的解释能力[4]; 4) 多模态需求: 推荐系统需处理文本、图像等多种数据, 增加了系统复杂性[2]。

其中最为关键的是计算成本与模型效果的权衡, 以及模型复杂性与可解释性的矛盾。本文针对这些挑战, 做出如下贡献:

1) 本文研究了通过合成坐标[36]将用户和物品映射到低维嵌入空间中, 利用空间关系和 KAN (Kolmogorov-Arnold Networks) [37]来学习用户与项目的相似性或偏好, 并以此来引导解释生成, 达到提升推荐解释生成的效果。

2) 本文设计了一种基于 Transformer 架构的, 融合合成坐标和 KAN 的引导式可解释推荐算法 GERT (Guided Explainable Recommendation Transformer), 该算法通过使用合成坐标来引导解释生成。

3) 本文评估了该算法在评分预测和解释生成任务中的表现。实验结果表明: 在评分预测任务中, 该算法达到了最先进的水平; 在解释生成任务中, 该算法在 BLEU-4、USR 等重要指标上有显著的提升。

4) 本文的研究揭示了, 通过优化嵌入向量的内容和模块化设计的结合, 可以有效提高可解释推荐系统的性能, 同时降低模型的复杂性、计算难度和资源需求。

2. 相关工作

近年来, 可解释推荐的发展呈现出趋势: 一方面大语言模型的广泛应用和自然语言增强与生成技术的突破, 使推荐解释更加流畅、自然; 另一方面则致力于通过多模态技术提升解释的全面性和准确性, 满足用户对推荐依据的深入理解。基于这一分析, 当前的研究主要围绕三种技术展开: 利用合成坐标进行空间向量建模、借助 KAN 捕捉非线性函数关系以及通过 Transformer 实现高质量解释文本的生成。

合成坐标(Synthetic Coordinates)最初是为了在高维数据中找到有效的表示形式, 以便在数据处理时减少维度和计算复杂度。研究表明[36][38], 通过将相关信息映射到一个合成的、多维度的坐标空间中, 利用空间关系(距离或方向)来衡量用户与项目的相似性或偏好, 可以提高推荐算法的效率和效果。在本文中, 合成坐标技术被用于将用户与物品映射到特征空间中, 并通过建模其空间关系实现评分预测。同时, 生成的空间向量也被用于引导后续的解释生成过程。

KAN [37]是一种基于 Kolmogorov-Arnold 定理的新型神经网络架构, 通过在网络的边缘(权重)上使用可学习的激活函数, 取代了传统多层感知器(MLP)中的固定激活函数。最新研究表明[37], KAN 能够在多维数据中捕捉复杂的非线性关系。此外, 将 KAN 与神经网络结合[39][40], 可以进一步提升模型的准确性。在本文中, KAN 被用于建模用户和物品的特征空间坐标与评分之间的非线性函数关系。

Transformer 是 Vaswani 等人于 2017 年提出的一种基于注意力机制的神经网络架构[41], 其初衷是为了解决序列建模任务中的长期依赖问题。Transformer 通过自注意力机制实现了并行计算, 显著提升了训练效率和效果。最新研究表明, Transformer 在多个领域表现出色[42]-[44], 已成为深度学习领域的核心, 其影响力和应用范围持续扩大[45]。在本文中, Transformer 被用于解释生成任务, 通过其强大的序列建模能力, 生成具有高质量和高可解释性的推荐解释文本。

在本文中, 合成坐标用于将用户与物品映射到特征空间, KAN 用于建模非线性空间关系以实现评分预测, Transformer 用于序列建模, 生成高质量的推荐解释文本。通过将合成坐标、KAN 和 Transformer 三种技术有机结合, 本文旨在设计并实现一种新型模型, 既能实现高效的评分预测, 又能生成符合用户

需求的高质量推荐解释，为可解释推荐系统的进一步发展提供理论与实践支持。

3. 方法

3.1. 问题定义

在可解释推荐的研究中，评分预测和解释生成是两个关键任务。给定一组用户和物品，评分预测主要关注通过模型预测用户对尚未接触的物品的评分，以此提供精准的推荐。预测的评分应当尽可能地接近真实评分。

另一方面，解释生成则旨在生成自然语言解释，使得推荐结果对用户更加透明和易于理解。生成的解释应具有连贯性和可读性，并且尽可能包含与用户和物品相关的特征集合。其中，解释文本和特征集合中的词汇元素均来自词汇表。

本文所用到的关键符号及其说明如表 1 所示。

Table 1. Key symbols and their descriptions
表 1. 关键符号及其说明

符号	说明
$\mathcal{T}$	训练集
$\mathcal{T}_e$	测试集
$\mathcal{U}$	用户集合
$\mathcal{I}$	物品集合
$\mathcal{V}$	词汇表
$\mathcal{F}$	特征集合
$\mathcal{E}$	解释集合
u, i	一组用户 - 物品对， u 代表用户， i 代表物品
s_u, s_i	用户 - 物品对 u 和 i 应的特征空间向量
$r_{u,i}$	用户 u 对物品 i 的真实评分
$c_{u,i}$	在整个词汇表中，与用户 u 对物品 i 相关词汇的概率分布
$F_{u,i}$	与用户 u 对物品 i 相关的特征集合
$E_{u,i}$	用户 u 对物品 i 的解释性文本序列
ψ, ϕ	简单函数参数
$\text{softmax}(\cdot)$	Softmax 函数
W	权重矩阵
b	偏置参数
$\mathcal{L}$	损失

3.2. 算法架构

如图 1 所示，整体结构可分为评分预测模块与解释生成模块。左侧为基于合成坐标和 KAN 的评分预测部分，右侧则是基于 Transformer 编码器 - 解码器架构的解释生成部分。其中，编码器(Encoder)负责上下文预测任务，相关研究表明[8] [9]，该任务的引入能够显著提升解释生成的效果。解码器(Decoder)则专注于解释生成任务的实现。

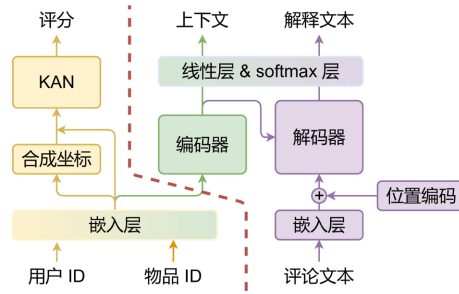


Figure 1. GERT model framework

图 1. GERT 的架构

在评分预测模块中，首先通过嵌入层(Embedding)将用户 id 和物品 id 编码为空间特征向量 s_u 和 s_i ，接着计算其距离 $d_{u,i}$ ，最终通过单层的 KAN 模型来建模其与真实评分 $r_{u,i}$ 的函数关系。

在上下文预测模块中，评分预测模块生成的空间特征向量 s_u 和 s_i 被用作编码器的输入。编码器输出的结果 x_c 经过单层线性变换和 Softmax 层，以估计解释中各单词的概率分布。

在解释生成模块中，主要采用自回归生成(Auto regressive Generation)方法来生成解释文本 $E_{u,i}$ 。具体来说，基于已有的单词序列 $e_1, e_2, \dots, e_t$ ，预测下一单词 e_{t+1} 。首先，序列 $e_1, e_2, \dots, e_t$ 经由嵌入层转换为对应的词向量，并通过位置编码进行位置标识，再与上下文预测模块生成的 x_c 一同输入到解码器。解码器输出 x_c 经由单层线性层和 Softmax 层，生成词汇表 V 中每个单词作为下一单词 e_t 的概率分布。以此方式逐词生成，直到整个句子的生成完成。

3.3. 基于合成坐标和 KAN 的评分预测

基于合成坐标的推荐系统[36][38]是一种创新的推荐方法，旨在通过对用户和项目(物品)之间的复杂关系进行建模和优化来提高推荐效果。在原论文中[36]，其建模方式如公式(1)所示：

$$d_{u,i} = \|s_u - s_i\|_2$$

$$\hat{r}_{u,i} = \max R - \frac{(\max R - \min R)}{100} d_{u,i} \quad (1)$$

其中， s_u 和 s_i 分别是合成的用户和物品的空间坐标， $d_{u,i}$ 是两坐标之间的距离， $\|\cdot\|_2$ 表示二范数，即欧几里得距离或向量的“长度”。 $\max R$ 和 $\min R$ 分别为数据集中的最大评分和最小评分。

然而，这种建模方式本质上是一种线性映射，本文认为这种简单的函数形式无法充分捕捉其中的复杂关系。因此，本文引入 KAN 来对更复杂的非线性关系进行建模，以提高模型的表现力和预测精度。

KAN 可以拟合任意多变量连续函数，可以将一个复杂的、 n 维多变量函数 $f(x_1, x_2, \dots, x_n)$ 分解成一系列更简单的单变量函数的组合，如公式(2)所示：

$$f(x_1, x_2, \dots, x_n) = \sum_{j=1}^{2n+1} \psi_j \left(\sum_{i=1}^n \phi_{i,j}(x_i) \right) \quad (2)$$

其中 x_i 是输入变量($i=1, 2, \dots, n$)， $\phi_{ij}(x_i)$ 和 ψ_j 都可以通过神经网络的方式来实现，从而有效地近似任意多变量函数。

本文使用合成的用户和物品的空间坐标 s_u 、 s_i ，及其空间关系 $d_{u,i}$ (距离)作为输入，并通过单层的 KAN 网络预测评分 $r_{u,i}$ 。使用的损失函数是均方根误差(MSE)，定义如公式(3)所示：

$$\hat{r}_{u,i} = \text{KAN}(s_u, s_i, d_{u,i})$$

$$\mathcal{L}_r = \frac{1}{|\mathcal{T}|} \sum_{(u,i) \in \mathcal{T}} (\hat{r}_{u,i} - r_{u,i})^2 \quad (3)$$

其中, $\mathcal{T}$ 是训练集数据, $||$ 表示集合中元素的数量。

伪代码(见算法 1)展示了该方法的训练过程。

算法 1. 基于合成坐标和 KAN 的评分预测算法的训练过程

Require 训练集 $T = (u, i, r(u, i))$, 验证集 $V = (u, i, r(u, i))$, 学习率 η , 训练轮数 N , 模型参数 θ
 Ensure 训练后的模型参数 θ^*

1. 初始化用户嵌入 $\mathbf{U}$ 和物品嵌入 $\mathbf{I}$ 为小的随机值
2. for $epoch \leftarrow 1$ to N
3. 随机打乱训练数据 T
4. for all $(u, i, r(u, i)) \in T$
5. $s_u \leftarrow \mathbf{U}[u]$ 获取用户嵌入
6. $s_i \leftarrow \mathbf{I}[i]$ 获取物品嵌入
7. $d_{u,i} \leftarrow \|s_u - s_i\|_2$ 计算欧几里得距离
8. $\hat{r}(u, i) \leftarrow \text{KAN}([s_u, s_i, d_{u,i}])$ 使用KAN模型预测评分
9. $\mathcal{L} \leftarrow \text{损失}(\hat{r}(u, i), r(u, i))$ 计算损失
10. 反向传播损失 $\mathcal{L}$ 并使用优化器更新参数 θ
11. end for
12. 初始化验证损失 $\mathcal{L}_{val} \leftarrow 0$
13. for all $(u, i, r(u, i)) \in D_{val}$
14. $s_u \leftarrow \mathbf{U}[u]$ 获取用户嵌入
15. $s_i \leftarrow \mathbf{I}[i]$ 获取物品嵌入
16. $d_{u,i} \leftarrow \|s_u - s_i\|_2$ 计算欧几里得距离
17. $\hat{r}(u, i) \leftarrow \text{KAN}([s_u, s_i, d_{u,i}])$ 使用KAN模型预测评分
18. $\mathcal{L}_{val} \leftarrow \mathcal{L}_{val} + \text{loss}(\hat{r}(u, i), r(u, i))$ 累计验证损失
19. end for
20. 记录并监控验证损失 $\mathcal{L}_{val}$
21. if $\mathcal{L}_{val}$ 没有改进(如连续几轮不减小)
22. 提前停止训练 避免过拟合
23. end if
24. end for
25. return θ^*

3.4. 基于 Transformer 的引导式解释生成

基于编码器的上下文预测任务中, 将合成的用户和物品的空间坐标 s_u 和 s_i 作为输入, 将其输出 x_c 通过一个 Softmax 层来计算解释中出现的所有单词的概率分布, 并通过交叉熵损失训练, 定义如公式(4)所示:

$$x_c = \text{Encoder}(s_u, s_i)$$

$$\hat{c}_{u,i} = \text{softmax}(W_c x_c + b_c)$$

$$\mathcal{L}_c = \frac{1}{|T|} \sum_{(u,i) \in T} \frac{1}{|E_{u,i}|} \sum_{t=1}^{|E_{u,i}|} -\log \hat{c}_{u,i}^{e_t} \quad (4)$$

其中, w_c 和 b_c 是可训练参数, e_t 是真实解释文本 $E_{u,i}$ 的第 t 个标记。

解码器部分主要负责生成最终的解释文本 $\hat{E}_{u,i}$, 使用交叉熵损失函数进行训练, 定义如公式(5)所示:

$$\begin{aligned} x_e &= \text{Decoder}(\hat{c}_{u,i}, E_{u,i}) \\ \hat{E}_{u,i} &= \text{softmax}(W_c x_e + b_c) \\ \mathcal{L}_e &= \frac{1}{|T|} \sum_{(u,i) \in T} \frac{1}{|E_{u,i}|} \sum_{t=1}^{|E_{u,i}|} -\log \hat{e}_t^{e_t} \end{aligned} \quad (5)$$

其中, w_e 和 b_e 是可训练参数, $\hat{e}_t$ 是生成解释文本 $\hat{E}_{u,i}$ 的第 t 个标记。

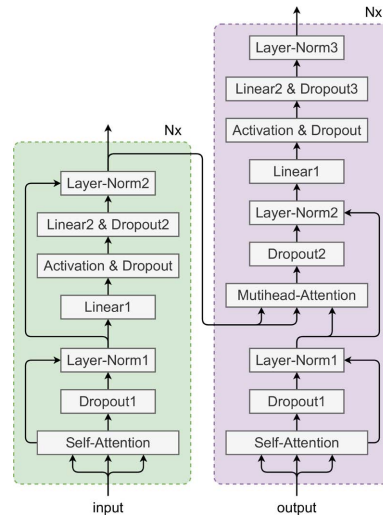


Figure 2. Details of the encoder (left)-decoder (right)
图 2. 编码器(左)-解码器(右)细节

编码器 - 解码器内部细节如图 2 所示。为描述解释生成的计算过程, 本文在算法 2 中展示了该方法的伪代码。

算法 2. 引导式解释生成

Require 用户特征 s_u , 物品特征 s_i , 文本 $[e_1, e_2, \dots, e_t]$, 模型参数 w_c , b_c , w_e , b_e

Ensure 上下文分布 $c_{u,i}$, 下一单词 e_t

1. $x_c \leftarrow \text{Encoder}([s_u, s_i])$ 编码器
2. $c_{u,i} \leftarrow \text{softmax}(W_c x_c + b_c)$
3. $src \leftarrow \text{Embedding}([e_1, e_2, \dots, e_t])$ 词嵌入
4. $src \leftarrow \text{PosEncoder}(src)$ 位置编码
5. $x_e \leftarrow \text{Decoder}(src, x_c)$ 解码器
6. $e_t \leftarrow \text{argmax}(\text{softmax}(W_e x_e + b_e))$
7. return $c_{u,i}, e_t$

3.5. 联合损失函数

最后，采用标准反向传播算法对整个模型进行训练。在此过程中，评分预测模块和解释生成模块既可以分别独立训练，也可以合并在一起进行联合训练。为此，定义联合损失函数为三个预测任务的加权和，如公式 6 所示：

$$L = \lambda_r \mathcal{L}_r + \lambda_c \mathcal{L}_c + \lambda_e \mathcal{L}_e \quad (6)$$

其中， λ_r 、 λ_c 和 λ_e 分别代表针对特定任务的权重系数，通过调整这些系数，可以在不同任务之间实现平衡，优化模型的整体性能。

4. 实验设计

4.1. 数据集与实验细节

本文使用三个公开的可解释推荐数据集，分别来自 Yelp (餐厅)、Amazon (电影和电视)和 TripAdvisor (酒店) [46]。在训练集中，每个用户和物品至少有一条记录，保证了数据的完整性和模型的训练效果。每条记录包含用户 ID、物品 ID、用户评分(1~5 星)、解释文本和特征信息。其中，解释文本是从用户评论中提取的句子，每个解释文本至少包含一个与推荐相关的特征。

Table 2. Statistics of the three datasets
表 2. 三个数据集的统计

	Yelp	Amazon	TripAdvisor
用户总数	27,147	7506	9765
物品总数	20,266	7360	6280
记录总数	1,293,247	441,783	320,023
特征总数	7340	5399	5069
用户平均记录数	47.64	58.86	32.77
物品平均记录数	63.81	60.02	50.96
解释文本平均长度	12.32	14.14	13.01
稀疏率(%)	99.76	99.20	99.48

表 2 显示了三个数据集的详细统计数据，可以从中看出各数据集在用户数、物品数、记录总数等方面的差异。例如，在数据稀疏性方面，Yelp 数据集用户和物品数量较多，但交互频率较低，稀疏性较高，适用于测试模型在高稀疏性条件下的表现。Amazon 数据集的交互更为集中，稀疏度较低。TripAdvisor 数据集介于两者之间，用户评分和解释文本集中于特定酒店物品，提供了一个相对均衡的测试环境。

每个数据集按照 8:1:1 的比例被随机划分为训练集、验证集和测试集。具体而言，训练集用于模型的学习和参数优化，验证集用于调整模型超参数和选择最佳模型，而测试集则用于评估模型在未见数据上的最终性能。为了确保结果的稳定性和可靠性，对每个数据集进行了 5 次重复实验。

在超参数的选择上，对嵌入向量维度在[128, 256, 384, 512, 768]范围内进行搜索，以确定最佳表现。评分预测任务的权重在[0.001, 0.01, 0.1, 1.0]中进行搜索，上下文预测任务和解释生成任务的权重均设定为 1.0。生成的句子长度设置为 15，学习率初始设为 1.0。模型在每个 epoch 达到最小损失时，学习率将调整为当前值的 0.25。当总计 3 个 epoch 的总损失没有显著下降时，训练将被终止。

4.2. 评价指标

为了评估可解释推荐的性能, 本论文采用 RMSE 和 MAE 来衡量评分预测的效果。MAE 是一种用于评估模型预测值与实际值之间误差的指标, 表示误差的绝对值平均大小。其计算方式如公式(7)所示:

$$MAE = \frac{1}{N} \sum_{u,i \in \mathcal{T}_e} |r_{u,i} - \hat{r}_{u,i}| \quad (7)$$

其中, N 为测试样本总数。

RMSE 也是一种用于评估模型预测值与实际值之间误差的指标, 但它比 MAE 更加关注较大的误差, 因为它对误差平方取平均后再开方。其计算方式如公式(8)所示:

$$RMSE = \sqrt{\frac{1}{N} \sum_{u,i \in \mathcal{T}_e} (r_{u,i} - \hat{r}_{u,i})^2} \quad (8)$$

在解释生成任务中, 本论文从可解释性和文本质量两个角度出发, 在可解释性评估中, 本论文采用特征匹配率(FMR)、特征覆盖率(FCR)和特征多样性(DIV) [46]。

FMR 被定义为生成的解释包含真实文本中的特征的比率, 计算方式如公式(9)所示:

$$FMR = \frac{1}{N} \sum_{u,i} \delta(f_{u,i} \in \hat{E}_{u,i})$$

$$\delta(x) = \begin{cases} 1 & \text{if } x = \text{true} \\ 0 & \text{if } x = \text{false} \end{cases} \quad (9)$$

其中, 对于每个用户-物品对 u,i , $\hat{E}_{u,i}$ 表示模型生成的解释文本, $f_{u,i}$ 表示真实特征。

FCR 被定义为所有生成的解释文本中包含的特征数量除以整个数据集中特征的数量, 计算方式如公式(10)所示:

$$FCR = \frac{|\hat{\mathcal{F}} \cap \mathcal{F}|}{|\mathcal{F}|} \quad (10)$$

其中, $\mathcal{F}$ 表示真实解释文本中特征的集合, $\hat{\mathcal{F}}$ 表示生成解释文本中特征的集合。

DIV 被定义为任意两个生成的文本之间的特征的交集, 其计算方式如公式(11)所示:

$$DIV = \frac{2}{N \times (N-1)} \sum_{u,i \in \mathcal{T}_e} |F_{u,i} \cap \hat{F}_{u',i'}| \quad (11)$$

其中, $\hat{F}_{u,i}$ 和 $\hat{F}_{u',i'}$ 表示两个生成的文本包含的特征集合。

在文本质量评估中, 本论文采用 USR [46]、BLEU-1、BLEU-4 [47]、ROUGE-1 和 ROUGE-2 [48] 的精确率(Precision)、召回率(Recall)、F1 分数来衡量解释生成的效果。

BLEU (Bilingual Evaluation Understudy) [47] 是用于评估生成句子质量的常用指标。其基本思想是通过计算生成句子与参考句子之间的 n-gram 匹配程度来评估翻译或生成任务的性能。BLEU 的计算过程主要包括以下几个步骤:

1) **计算 n-gram 精确度(Precision):** 对于每个 n-gram, 计算生成句子 $\hat{E}$ 中每个 n-gram 与参考句子 E 中 n-gram 的匹配个数。然后, 计算生成句子中所有 n-gram 的精确度。

2) **计算 BP 惩罚因子:** 为了防止生成的句子过短而导致 n-gram 精确度较高, BLEU 引入了惩罚因子 BP (Brevity Penalty), 其计算方式如公式(12)所示:

$$BP = \begin{cases} 1 & \text{if } \frac{|\hat{E}|}{|E|} > 1 \\ e^{\left(1 - \frac{|\hat{E}|}{|E|}\right)} & \text{if } \frac{|\hat{E}|}{|E|} \leq 1 \end{cases} \quad (12)$$

3) **计算 BLEU 分数:** 结合不同 n-gram 的精确度和 BP 惩罚因子, BLEU 的计算方式如公式(13)所示:

$$BLEU = BP \cdot \exp\left(\sum_{n=1}^N w_n \log P_n\right) \quad (13)$$

其中, P_n 表示第 n 阶 n-gram 的精确度, w_n 为权重, 通常取 $w_n = 1/N$ 。

BLEU-1 和 BLEU-4 分别代表只考虑 1-gram 精确度和考虑 1-gram 到 4-gram 精确度的 BLEU 得分, 其计算方式如公式(14)所示:

$$\begin{aligned} BLEU-1 &= BP \cdot P_1 \\ BLEU-4 &= BP \cdot \exp\left(\frac{1}{4}(\log P_1 + \log P_2 + \log P_3 + \log P_4)\right) \end{aligned} \quad (14)$$

BLEU 分数越高, 表示生成句子 $\hat{E}$ 与参考句子 E 之间的相似度越高, 生成质量也越好。

ROUGE (Recall-Oriented Understudy for Gisting Evaluation) [48]是文本摘要任务中常用的评价指标, 用于衡量生成的句子 $\hat{E}$ 与原始句子 E 之间的相似度。ROUGE 的主要度量包括精确度(Precision)、召回率(Recall)和 F1 分数。它们的计算过程如下:

1) **精确度(Precision):** 精确度衡量生成的句子 $\hat{E}$ 中与原始句子 E 相匹配的部分占 $\hat{E}$ 总长度的比例, 计算方式如公式(15)所示:

$$Precision = \frac{|\hat{E} \cap E|}{|\hat{E}|} \quad (15)$$

2) **召回率(Recall):** 召回率衡量原始句子 E 中与生成句子 $\hat{E}$ 相匹配的部分占 E 总长度的比例, 计算方式如公式(16)所示:

$$Recall = \frac{|\hat{E} \cap E|}{|E|} \quad (16)$$

3) **F1 分数(F1-Score):** F1 分数是精确度和召回率的调和平均, 用于综合评估这两者的平衡性, 计算方式如公式(17)所示:

$$F1-Score = \frac{2 \cdot Precision \cdot Recall}{Precision + Recall} \quad (17)$$

其中, $|\hat{E} \cap E|$ 表示生成的句子 $\hat{E}$ 和原始句子 E 之间匹配的词语或片段数量, $|\hat{E}|$ 和 $|E|$ 分别表示生成句子和原始句子的长度(通常以词或 n-gram 的数量来衡量)。

然而, BLEU 和 ROUGE 主要通过词汇匹配来评估生成文本的质量, 很难检测相同句子, 因此在评估生成文本的个性化程度方面存在局限性。为了更好地评估个性化生成文本, 本论文采用 USR 计算生成唯一性句子的比率, 计算方式如公式(18)所示:

$$USR = \frac{|\hat{\mathcal{E}}|}{N} \quad (18)$$

其中 $\hat{\mathcal{E}}$ 表示模型生成的句子集合。

研究表明用户对机器生成解释的感知与其相关性、重复性以及特征外观等因素密切相关[49], 这些因素分别对应评价指标中的 BLEU/ROUGE、USR 和 FMR。

4.3. 基线算法

本论文考虑两组基线, 分别对评分预测和解释生成任务进行自动评估。在评分预测方面, 本论文使用两种基于分解的方法 PMF [50]、SVD++ [51]和两种基于神经网络的方法 NRT [52]和 PETER [8]作为比较基准。在解释生成方面, 采用 NRT、Att2Seq [53]、PETER 和 PEPLER [27]。以下是几种基线算法的介绍。

PMF: 一种基于概率的矩阵分解方法, 用于协同过滤推荐系统。它通过矩阵分解将用户 - 项目评分矩阵分解为两个低维矩阵, 分别表示用户和项目的特征向量, 通过矩阵乘积预测未知评分。

SVD++: 是对经典 SVD (奇异值分解)方法的扩展。它在基本的矩阵分解之外, 还考虑了用户对项目的隐式反馈(例如用户点击或浏览记录), 从而增强了对用户偏好的捕捉。

NRT: 一种神经网络模型, 用于生成评分和提示。它结合了深度学习技术, 通过神经网络对用户和项目的交互进行建模, 从而生成个性化的评分预测和相关提示。

Att2Seq: 一种从属性生成序列的模型, 通常用于生成个性化推荐的文本解释。它使用了序列到序列(Seq2Seq)的框架, 能够根据输入的属性生成相关的输出序列, 适用于推荐解释生成任务。

PETER: 利用了 Transformer 模型的强大表示能力, 将评分预测任务与解释生成任务结合在一起。通过多任务学习, 模型能够在提供准确评分预测的同时, 生成更具个性化的解释。

PEPLER: 基于 GPT-2 [16]的模型, 旨在通过个性化提示学习[54]提供可解释的推荐。它通过微调 GPT-2 模型, 使其能够生成个性化的推荐解释。

其中, PETER 的实验数据由本文实验结果获得, 其参数设置采用了论文[8]中的最优配置。由于 PEPLER 的计算时间过长且计算资源有限, 本文仅进行了单次实验。实验结果表明, PETER 和 PEPLER 的性能与论文[8][9][27]中的表现基本一致, 验证了其实验的准确性。因此, 对于部分基线算法的实验数据, 本文直接引用了论文[8][9][27]中的结果, 以确保研究的完整性和可靠性。

5. 实验结果分析

5.1. 评分预测

评分预测结果如表 3 和图 3 所示。可以看出, GERT 在评分预测任务中展示出较强的竞争力, 尤其是在 TripAdvisor 数据集上表现最佳。TripAdvisor 数据集的用户和物品平均记录数相对较低, 这可能使得模型在学习用户偏好时更加困难。但 GERT 在该数据集上取得了最优的 RMSE (0.79)和 MAE (0.60)结果, 表明其在处理该类型的数据时表现尤为突出。

然而, 在 Yelp 和 Amazon 数据集上, GERT 模型的表现略逊于 NRT 和 PETER, 这表明其在这两个数据集上可能未能达到最佳的预测精度。这种局限性可能源于 GERT 模型在处理某些特定数据特征或特定推荐任务时的不足。具体而言, GERT 可能在应对数据稀疏性、用户行为模式的复杂性, 或在模型参数调优方面存在一定的挑战。

5.2. 解释生成

在解释生成任务中, 各算法的表现已在表 4 进行了详细展示。就可解释性指标而言, 各算法在 FMR (解释准确性率)指标上的表现较为接近。然而, 在 FCR (解释覆盖率)和 DIV (解释多样性)指标上, 结果却

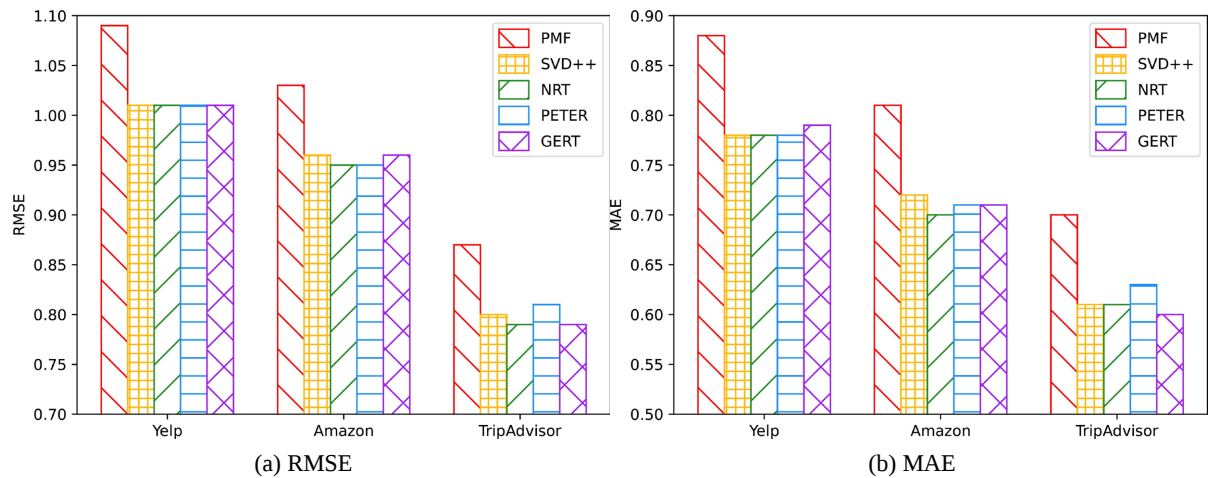


Figure 3. Rating prediction results: RMSE (a) and MAE (b)

图 3. 评分预测结果: RMSE (a)和 MAE (b)

Table 3. Rating prediction results, with the best performance values in bold

表 3. 评分预测结果, 最佳性能值以粗体标出

	Yelp		Amazon		TripAdvisor	
	RMSE	MAE	RMSE	MAE	RMSE	MAE
PMF	1.09	0.88	1.03	0.81	0.87	0.70
SVD++	1.01	0.78	0.96	0.72	0.80	0.61
NRT	1.01	0.78	0.95	0.70	0.79	0.61
PETER	1.01	0.78	0.95	0.71	0.81	0.63
GERT	1.01	0.79	0.96	0.71	0.79	0.60

有所不同, 具体情况见图 4(a)和图 4(b)所示。尽管 GERT 在 FCR 值上略逊于 PEPLER, 但其总体表现仍优于其他算法。特别值得注意的是, GERT 在 DIV 值上表现尤为突出, 尤其是在 Amazon 和 TripAdvisor 数据集上, 显示了其在生成多样化解释文本方面的显著优势。

在文本质量方面, 各算法在 USR 和 BLEU-4 上的表现如图 4(c)和图 4(d)所示。GERT 在这两个指标上均取得了优异的成绩, 这表明 GERT 在文本生成任务中表现出强大的能力, 能够生成既高质量又符合用户需求的解释文本。这一结果进一步验证了 GERT 在提升解释文本质量和满足用户期望方面的有效性, 突显了其在生成内容的准确性和用户满意度方面的显著优势。

5.3. 超参数与消融实验

本论文旨在研究合成坐标引导的解释生成效果及其影响。为实现这一目标, 本文通过详细的实验分析了联合损失函数中正则化系数 λ_r (超参数) 的不同取值对模型表现的具体影响。同时, 本文还探索了 KAN 在捕捉非线性函数特征方面的作用, 为此进行了消融实验, 并采用了单层多层感知器(MLP)进行对比分析。此外, 为了更全面地评估和比较模型的效果, 本文还将 PETER 模型纳入对比范围。所有实验均在 TripAdvisor 数据集上进行, 其中 GERT 模型的嵌入层大小设定为 128。

实验结果如图 5 所示, 结果表明, λ_r 取值对解释生成的精度和稳定性具有显著的影响。当正则化系数 $\lambda_r = 10^{-2}$ 时, 模型的表现达到了最佳状态。此外, 该实验进一步验证了 KAN 在捕捉复杂非线性函数关系方面的有效性。与传统的单层 MLP 相比, KAN 展现出了更加优越的性能, 能够更好地处理和捕捉数据中的复杂特征。

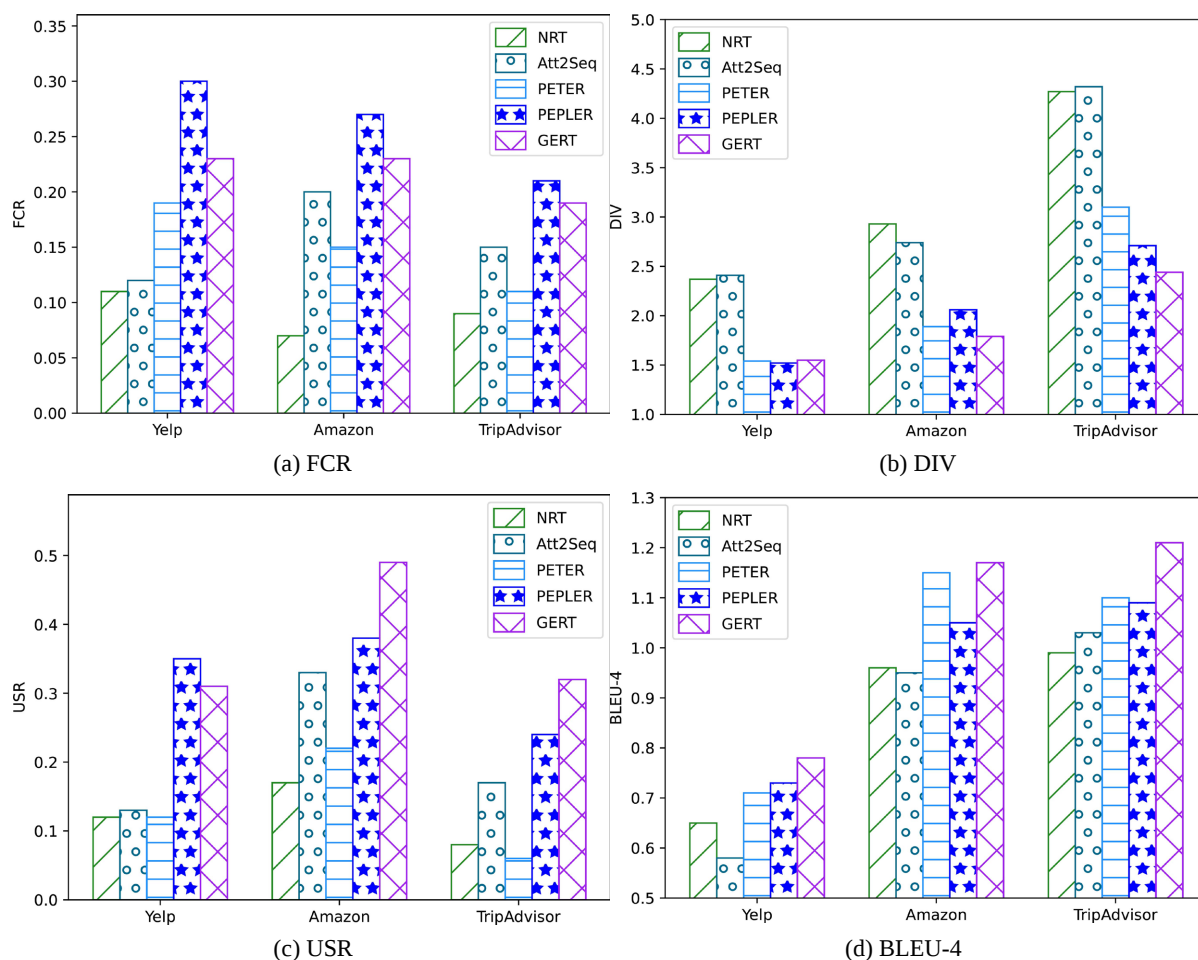


Figure 4. 解释生成结果: FCR (a); DIV (b); USR (c)和 BLEU-4 (d)

图 4. Explanation generation results: FCR (a); DIV (b); USR (c) and BLEU-4 (d)

Table 4. Comparison of text quality, with the best performance values in bold. B1 and B4 represent BLEU-1 and BLEU-4. R1-P, R1-R, R1-F, R2-P, R2-R, and R2-F represent the precision, recall, and F1 value of ROUGE-1 and ROUGE-2. The results of BLEU and ROUGE are expressed in percentages (% omitted), and the other indicators are expressed in absolute values

表 4. 解释生成结果。评估指标 B1 和 B4 分别代表 BLEU-1 和 BLEU-4。R1-P、R1-R、R1-F、R2-P、R2-R 和 R2-F 分别表示 ROUGE-1 和 ROUGE-2 的精度(precision)、召回率(recall)和 F1 值。BLEU 和 ROUGE 的结果以百分比表示(省略%), 其他指标以绝对值表示。最佳性能值用粗体标出

	可解释性				文本质量							
	FMR	FCR	DIV ↓	USR	B1	B4	R1-P	R1-R	R1-F	R2-P	R2-R	R2-F
Yelp												
NRT	0.07	0.11	2.37	0.12	11.66	0.65	17.69	12.11	13.55	1.76	1.22	1.33
Att2Seq	0.07	0.12	2.41	0.13	10.29	0.58	18.73	11.28	13.29	1.85	1.14	1.31
PETER	0.08	0.19	1.54	0.12	10.64	0.71	18.57	12.10	13.70	2.02	1.35	1.47
PEPLER	0.08	0.30	1.52	0.35	11.23	0.73	17.51	12.55	13.53	1.86	1.42	1.46
GERT	0.08	0.23	1.55	0.31	11.43	0.78	17.59	12.46	13.66	1.92	1.47	1.52
Amazon												
NRT	0.12	0.07	2.93	0.17	12.93	0.96	21.03	13.57	15.56	2.71	1.84	2.05

续表

Att2Seq	0.12	0.20	2.74	0.33	12.56	0.95	20.79	13.31	15.35	2.62	1.78	1.99
PETER	0.12	0.15	1.89	0.22	12.97	1.15	20.08	13.95	15.43	2.82	2.10	2.22
PEPLER	0.11	0.27	2.06	0.38	13.19	1.05	18.51	14.16	14.87	2.36	1.88	1.91
GERT	0.12	0.23	1.79	0.49	13.34	1.17	19.33	14.11	15.28	2.65	2.07	2.14

TripAdvisor												
NRT	0.06	0.09	4.27	0.08	15.05	0.99	18.22	14.39	15.4	2.29	1.98	2.01
Att2Seq	0.06	0.15	4.32	0.17	15.27	1.03	18.97	14.72	15.92	2.40	2.03	2.09
PETER	0.07	0.11	3.10	0.06	15.98	1.10	18.90	16.02	16.37	2.33	2.17	2.10
PEPLER	0.07	0.21	2.71	0.24	15.49	1.09	19.48	15.67	16.24	2.48	2.21	2.16
GERT	0.07	0.19	2.44	0.32	15.96	1.21	19.39	15.98	16.56	2.44	2.24	2.18

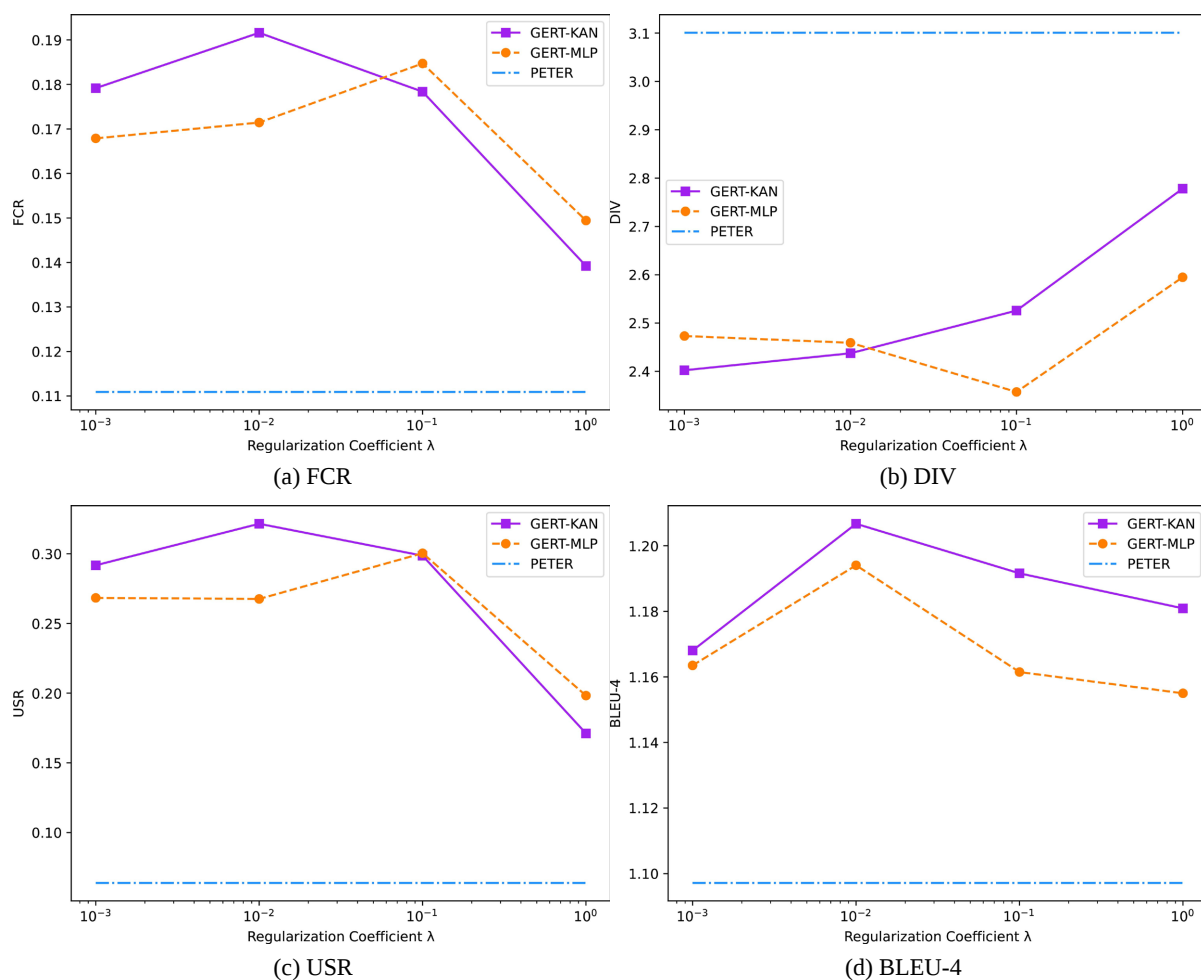


Figure 5. Recommended task weight and ablation experiment

图 5. 推荐任务权重与消融实验

5.4. 结果分析

在数据相对密集的情境下(例如在 Amazon 和 TripAdvisor 数据集上), GERT 算法展现了卓越的性能。

这一表现主要归因于其通过合成坐标嵌入技术，能够有效地提取用户和物品的潜在特征，并利用 KAN 模型深入挖掘这些特征之间的复杂关系。这种嵌入不仅显著提高了推荐的准确性，还为生成解释性文本提供了丰富的上下文支持，使得生成的文本在相关性和说服力方面表现更加出色。

然而，在面对数据相对稀疏的数据集时(例如 Yelp)，GERT 算法的表现却暴露出了一定的局限性。这些局限性可能源于以下几个方面：

首先，数据稀疏性可能导致模型在特征提取和嵌入过程中无法充分捕捉用户与物品之间的关系，进而影响推荐精度和生成文本的相关性。

其次，在生成过程中，由于缺乏足够的上下文信息，GERT 算法的引导机制可能无法完全避免局部暗点问题，导致生成的内容未能充分反映用户的偏好。

此外，稀疏数据集中的噪声和数据不平衡问题可能进一步削弱模型的表现，使得生成的推荐结果和解释性文本缺乏稳定性和连贯性。

最后，GERT 算法可能在特征提取和关系建模上高度依赖充足的数据支持，而稀疏数据集难以提供足够的信息，从而限制了算法优势的充分发挥。

6. 结论

本文提出了一种结合合成坐标和 KAN 的引导式可解释推荐方法。该方法通过将用户和物品映射到嵌入空间中，利用空间关系引导解释生成，从而提升推荐系统的可解释性。

在模型性能方面，实验结果表明：1) 该方法在评分预测任务中展现出强大的竞争力，达到了当前先进水平。2) 在解释生成任务中，GERT 在 DIV、USR、BLEU-4 等指标上表现出色，证明了其在高质量文本生成方面的卓越能力。

在模型复杂性方面，该方法通过将可解释推荐任务有针对性地分解为评分预测和解释生成两个独立模块，实现了更高的灵活性，降低了内部耦合度。此外，利用空间关系将两个模块进行有效连接，提升了系统在算法层面的可解释性。

在计算难度方面，该方法允许独立训练评分预测和解释生成模块，从而降低了整体训练的难度和资源需求。同时，与 PETER 和 PEPLER 相比，GERT 拥有更低的嵌入维度。例如，在 TripAdvisor 数据集上，GERT 展现出更优异的生成效果，但其嵌入维度仅为 128，远低于 PETER 的 512 维和 PEPLER 的 768 维。这不仅意味着 GERT 拥有更低的模型复杂度，更高的信息密度，同时也提升了计算效率和存储效率。

本文的研究揭示了，通过优化嵌入向量的内容和模块化设计的结合，可以有效提高可解释推荐系统的性能，同时降低模型的复杂性、计算难度和资源需求。该方法证明了其在提升推荐系统性能和可解释性方面的潜力，并为构建更高效的推荐系统提供了新的思路和方法。

未来，随着推荐系统在各类应用中的广泛使用，对可解释性与高效性的需求将持续增长。基于本研究提出的融合合成坐标和 KAN 的引导式可解释推荐方法，未来研究可以进一步探讨以下几个方向：

首先，在模型的可扩展性和适应性方面，未来可以尝试将该方法应用于更加复杂和多样化的数据集，以验证其在不同场景中的通用性和鲁棒性。

其次，可以进一步优化嵌入空间的构建与空间关系的表达方式，以提升模型在处理高维数据和复杂用户行为模式时的精度和效率。

此外，随着深度学习技术的不断发展，未来的研究还可以探索将其他先进的神经网络结构与该方法相结合，以进一步提高推荐系统的可解释性和性能。

基金项目

国家自然科学基金项目(61803264)。

参考文献

- [1] Deldjoo, Y., He, Z., McAuley, J., Korikov, A., Sanner, S., Ramisa, A., *et al.* (2024) A Review of Modern Recommender Systems Using Generative Models (Gen-RecSys). *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, Barcelona, 25-29 August 2024, 6448-6458. <https://doi.org/10.1145/3637528.3671474>
- [2] Chatti, M.A., Guesmi, M. and Muslim, A. (2024) Visualization for Recommendation Explainability: A Survey and New Perspectives. *ACM Transactions on Interactive Intelligent Systems*, **14**, 1-40. <https://doi.org/10.1145/3672276>
- [3] Zhang, Y. and Chen, X. (2020) Explainable Recommendation: A Survey and New Perspectives. *Foundations and Trends® in Information Retrieval*, **14**, 1-101. <https://doi.org/10.1561/15000000066>
- [4] Ma, B., Yang, T. and Ren, B. (2024) A Survey on Explainable Course Recommendation Systems. In: Streitz, N.A. and Konomi, S., Eds., *Distributed, Ambient and Pervasive Interactions*, Springer, 273-287. https://doi.org/10.1007/978-3-031-60012-8_17
- [5] Becker, J., Wahle, J.P., Gipp, B., *et al.* (2024) Text Generation: A Systematic Literature Review of Tasks, Evaluation, and Challenges. <https://arxiv.org/abs/2405.15604>
- [6] Lee, Y., Ka, S., Son, B., *et al.* (2024) Navigating the Path of Writing: Outline-Guided Text Generation with Large Language Models. <https://arxiv.org/abs/2404.13919>
- [7] Lin, Z., Guan, S., Zhang, W., Zhang, H., Li, Y. and Zhang, H. (2024) Towards Trustworthy LLMs: A Review on Debiasing and Dehallucinating in Large Language Models. *Artificial Intelligence Review*, **57**, Article No. 243. <https://doi.org/10.1007/s10462-024-10896-y>
- [8] Li, L., Zhang, Y. and Chen, L. (2021) Personalized Transformer for Explainable Recommendation. *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing*, Volume 1, 4947-4957. <https://doi.org/10.18653/v1/2021.acl-long.383>
- [9] Raczyński, J., Lango, M. and Stefanowski, J. (2023) The Problem of Coherence in Natural Language Explanations of Recommendations. In: Gal, K., Nowé, A., Nalepa, G.J., *et al.*, Eds., *Frontiers in Artificial Intelligence and Applications*, IOS Press, 1922-1929. <https://doi.org/10.3233/faia230482>
- [10] Liu, Z., Ma, Y., Schubert, M., Ouyang, Y., Rong, W. and Xiong, Z. (2024) Multimodal Contrastive Transformer for Explainable Recommendation. *IEEE Transactions on Computational Social Systems*, **11**, 2632-2643. <https://doi.org/10.1109/tcss.2023.3276273>
- [11] Wang, L., Cai, Z., De Melo, G., Cao, Z. and He, L. (2023) Disentangled CVAEs with Contrastive Learning for Explainable Recommendation. *Proceedings of the AAAI Conference on Artificial Intelligence*, **37**, 13691-13699. <https://doi.org/10.1609/aaai.v37i11.26604>
- [12] He, M., An, B., Wang, J. and Wen, H. (2023) CETD: Counterfactual Explanations by Considering Temporal Dependencies in Sequential Recommendation. *Applied Sciences*, **13**, Article No. 11176. <https://doi.org/10.3390/app132011176>
- [13] Guo, Y., Cai, F., Chen, H., Chen, C., Zhang, X. and Zhang, M. (2023) An Explainable Recommendation Method Based on Diffusion Model. *2023 9th International Conference on Big Data and Information Analytics (BigDIA)*, Haikou, 15-17 December 2023, 802-806. <https://doi.org/10.1109/bigdia60676.2023.10429319>
- [14] Li, L., Li, S., Marantika, W., *et al.* (2024) Diffusion-EXR: Controllable Review Generation for Explainable Recommendation via Diffusion Models. <https://arxiv.org/abs/2312.15490>
- [15] Minaee, S., Mikolov, T., Nikzad, N., *et al.* (2024) Large Language Models: A Survey. <https://arxiv.org/abs/2402.06196>
- [16] Radford, A., Wu, J., Child, R., *et al.* (2019) Language Models Are Unsupervised Multitask Learners. <https://www.semanticscholar.org/paper/Language-Models-are-Unsupervised-Multitask-Learners-Radford-Wu/9405cc0d6169988371b2755e573cc28650d14dfe>
- [17] Brown, T.B., Mann, B., Ryder, N., *et al.* (2020) Language Models Are Few-Shot Learners. In: *Proceedings of the 34th International Conference on Neural Information Processing Systems*, Curran Associates Inc., 1-25.
- [18] OpenAI, Achiam, J., Adler, S., *et al.* (2024) GPT-4 Technical Report. <https://arxiv.org/abs/2303.08774>
- [19] Devlin, J., Chang, M., Lee, K. and Toutanova, K. (2019) BERT: Pre-Training of Deep Bidirectional Transformers for Language Understanding. In: Burstein, J., Doran, C. and Solorio, T., Eds., *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, Association for Computational Linguistics, 4171-4186. <https://doi.org/10.18653/v1/n19-1423>
- [20] Luo, Y., Cheng, M., Zhang, H., Lu, J. and Chen, E. (2024) Unlocking the Potential of Large Language Models for Explainable Recommendations. In: Onizuka, M., *et al.*, Eds., *Database Systems for Advanced Applications*, Springer, 286-303. https://doi.org/10.1007/978-981-97-5569-1_18
- [21] Chu, Z., Wang, Y., Cui, Q., *et al.* (2024) LLM-Guided Multi-View Hypergraph Learning for Human-Centric Explainable Recommendation. <https://arxiv.org/abs/2401.08217>

- [22] Ma, Q., Ren, X. and Huang, C. (2024) XRec: Large Language Models for Explainable Recommendation. *Findings of the Association for Computational Linguistics: EMNLP 2024*, Miami, November 2024, 391-402. <https://doi.org/10.18653/v1/2024.findings-emnlp.22>
- [23] Yang, M., Zhu, M., Wang, Y., Chen, L., Zhao, Y., Wang, X., et al. (2024) Fine-Tuning Large Language Model Based Explainable Recommendation with Explainable Quality Reward. *Proceedings of the AAAI Conference on Artificial Intelligence*, **38**, 9250-9259. <https://doi.org/10.1609/aaai.v38i8.28777>
- [24] Lin, C., Tsai, C., Su, S., Jwo, J., Lee, C. and Wang, X. (2023) Predictive Prompts with Joint Training of Large Language Models for Explainable Recommendation. *Mathematics*, **11**, Article No. 4230. <https://doi.org/10.3390/math11204230>
- [25] Peng, Q., Liu, H., Xu, H., et al. (2024) Review-LLM: Harnessing Large Language Models for Personalized Review Generation. <https://arxiv.org/abs/2407.07487>
- [26] Peng, Y., Chen, H., Lin, C., et al. (2024) Uncertainty-Aware Explainable Recommendation with Large Language Models. <https://arxiv.org/abs/2402.03366>
- [27] Li, L., Zhang, Y. and Chen, L. (2023) Personalized Prompt Learning for Explainable Recommendation. *ACM Transactions on Information Systems*, **41**, 1-26. <https://doi.org/10.1145/3580488>
- [28] Lyu, Z., Wu, Y., Lai, J., Yang, M., Li, C. and Zhou, W. (2022) Knowledge Enhanced Graph Neural Networks for Explainable Recommendation. *IEEE Transactions on Knowledge and Data Engineering*, **35**, 4954-4968. <https://doi.org/10.1109/tkde.2022.3142260>
- [29] Bastola, R. and Shakya, S. (2024) Knowledge-Enriched Graph Convolution Network for Hybrid Explainable Recommendation from Review Texts and Reasoning Path. *2024 International Conference on Inventive Computation Technologies (ICICT)*, Lalitpur, 24-26 April 2024, 590-599. <https://doi.org/10.1109/icit60155.2024.10544384>
- [30] Qu, X., Wang, Y., Li, Z. and Gao, J. (2024) Graph-Enhanced Prompt Learning for Personalized Review Generation. *Data Science and Engineering*, **9**, 309-324. <https://doi.org/10.1007/s41019-024-00252-z>
- [31] Li, J., He, Z., Shang, J. and McAuley, J. (2023) Uceplic: Unifying Aspect Planning and Lexical Constraints for Generating Explanations in Recommendation. *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, Long Beach, 6-10 August 2023, 1248-1257. <https://doi.org/10.1145/3580305.3599535>
- [32] Cheng, H., Wang, S., Lu, W., Zhang, W., Zhou, M., Lu, K., et al. (2023) Explainable Recommendation with Personalized Review Retrieval and Aspect Learning. *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics*, Volume 1, 51-64. <https://doi.org/10.18653/v1/2023.acl-long.4>
- [33] Hu, Y., Liu, Y., Miao, C., Lin, G. and Miao, Y. (2022) Aspect-Guided Syntax Graph Learning for Explainable Recommendation. *IEEE Transactions on Knowledge and Data Engineering*, **35**, 7768-7778. <https://doi.org/10.1109/tkde.2022.3221847>
- [34] Liao, H., Wang, S., Cheng, H., Zhang, W., Zhang, J., Zhou, M., et al. (2025) Aspect-Enhanced Explainable Recommendation with Multi-Modal Contrastive Learning. *ACM Transactions on Intelligent Systems and Technology*, **16**, 1-24. <https://doi.org/10.1145/3673234>
- [35] Yan, A., He, Z., Li, J., Zhang, T. and McAuley, J. (2023) Personalized Showcases: Generating Multi-Modal Explanations for Recommendations. *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*, Taipei City, 23-27 July 2023, 2251-2255. <https://doi.org/10.1145/3539618.3592036>
- [36] Papadakis, H., Panagiotakis, C. and Fragopoulou, P. (2017) Scor: A Synthetic Coordinate Based Recommender System. *Expert Systems with Applications*, **79**, 8-19. <https://doi.org/10.1016/j.eswa.2017.02.025>
- [37] Liu, Z., Wang, Y., Vaidya, S., et al. (2024) KAN: Kolmogorov-Arnold Networks. <https://arxiv.org/abs/2404.19756>
- [38] Panagiotakis, C., Papadakis, H. and Fragopoulou, P. (2020) A User Training Error Based Correction Approach Combined with the Synthetic Coordinate Recommender System. *Adjunct Publication of the 28th ACM Conference on User Modeling, Adaptation and Personalization*, Genoa, 12-18 July 2020, 11-16. <https://doi.org/10.1145/3386392.3397591>
- [39] Bodner, A.D., Tepsich, A.S., Spolski, J.N., et al. (2024) Convolutional Kolmogorov-Arnold Networks. <https://arxiv.org/abs/2406.13155>
- [40] Kiamari, M., Kiamari, M. and Krishnamachari, B. (2024) GKAN: Graph Kolmogorov-Arnold Networks. <https://arxiv.org/abs/2406.06470>
- [41] Vaswani, A., Shazeer, N., Parmar, N., et al. (2017) Attention Is All You Need. *Proceedings of the 31st International Conference on Neural Information Processing Systems*, Long Beach, 4-9 December 2017, 6000-6010.
- [42] Cooper, J.F., Choi, S.J. and Buyuktaktakin, I.E. (2024) Toward Transformers: Revolutionizing the Solution of Mixed Integer Programs with Transformers. <https://arxiv.org/abs/2402.13380>
- [43] Ahn, K., Cheng, X., Song, M., et al. (2024) Linear Attention Is (Maybe) All You Need (to Understand Transformer Optimization). <https://arxiv.org/abs/2310.01082>

-
- [44] Hatamizadeh, A., Song, J., Liu, G., Kautz, J. and Vahdat, A. (2024) DiffiT: Diffusion Vision Transformers for Image Generation. *Computer Vision-ECCV 2024*, Milan, 29 September-4 October 2024, 37-55. https://doi.org/10.1007/978-3-031-73242-3_3
 - [45] Amatriain, X., Sankar, A., Bing, J., *et al.* (2024) Transformer Models: An Introduction and Catalog. <https://arxiv.org/abs/2302.07730>
 - [46] Li, L., Zhang, Y. and Chen, L. (2020) Generate Neural Template Explanations for Recommendation. *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*, 19-23 October 2020, 755-764. <https://doi.org/10.1145/3340531.3411992>
 - [47] Papineni, K., Roukos, S., Ward, T. and Zhu, W. (2001) BLEU: A Method for Automatic Evaluation of Machine Translation. *Proceedings of the 40th Annual Meeting on Association for Computational Linguistics-ACL'02*, Philadelphia, 6-12 July 2002, 311-318. <https://doi.org/10.3115/1073083.1073135>
 - [48] Lin, C.Y. (2004) Rouge: A Package for Automatic Evaluation of Summaries. In: *Text Summarization Branches Out*, Association for Computational Linguistics, 74-81. <https://aclanthology.org/W04-1013>
 - [49] Wen, B., Feng, Y., Zhang, Y. and Shah, C. (2022) Expscore: Learning Metrics for Recommendation Explanation. *Proceedings of the ACM Web Conference 2022*, Lyon, 25-29 April 2022, 3740-3744. <https://doi.org/10.1145/3485447.3512269>
 - [50] Mnih, A. and Salakhutdinov, R.R. (2007) Probabilistic Matrix Factorization. In: Platt, J., Koller, D., Singer, Y., *et al.*, Eds., *Advances in Neural Information Processing Systems*, Vol. 20, Curran Associates, Inc., 1257-1264. https://proceedings.neurips.cc/paper_files/paper/2007/file/d7322ed717dedf1eb4e6e52a37ea7bcd-Paper.pdf
 - [51] Koren, Y. (2008) Factorization Meets the Neighborhood: A Multifaceted Collaborative Filtering Model. *Proceedings of the 14th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, Las Vegas, 24-27 August 2008, 426-434. <https://doi.org/10.1145/1401890.1401944>
 - [52] Li, P., Wang, Z., Ren, Z., Bing, L. and Lam, W. (2017) Neural Rating Regression with Abstractive Tips Generation for Recommendation. *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval*, Tokyo, 7-11 August 2017, 345-354. <https://doi.org/10.1145/3077136.3080822>
 - [53] Dong, L., Huang, S., Wei, F., Lapata, M., Zhou, M. and Xu, K. (2017) Learning to Generate Product Reviews from Attributes. *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics*, 1, 623-632. <https://doi.org/10.18653/v1/e17-1059>
 - [54] Liu, P., Yuan, W., Fu, J., Jiang, Z., Hayashi, H. and Neubig, G. (2023) Pre-Train, Prompt, and Predict: A Systematic Survey of Prompting Methods in Natural Language Processing. *ACM Computing Surveys*, 55, 1-35. <https://doi.org/10.1145/3560815>