

序列推荐启发式数据增强的修复和丰富^{*}

党翌洲, 赵 楚, 姜琳颖, 马连博, 郭贵冰

(东北大学 软件学院, 辽宁 沈阳 110169)

通信作者: 郭贵冰, E-mail: guogb@swc.neu.edu.cn



摘 要: 启发式数据增强通过扰动原始的用户行为序列并生成新的序列, 能够有效缓解序列推荐中的数据稀疏性. 相比于基于外部信息、利用复杂的增强规则或引入可学习生成模块的增强方法, 启发式方法通过基于人类先验知识、规则或直觉设计的简单操作来完成增强, 具备较高的效率及泛化性. 然而, 现有的启发式数据增强通常面临两个问题: 一是其产生的增强数据与原始数据相比存在语义偏差或重要用户行为丢失的问题, 导致模型学习到不准确甚至是错误的用户偏好; 二是现有方法专注于对稀疏的输入行为进行增强, 模型训练仍受限于标签行为的稀疏性和不匹配问题, 导致推荐性能提升受限. 为此, 提出一种用于修复和丰富序列推荐中启发式数据增强的框架 (REDRec). 具体而言, 首先通过实证研究验证启发式推荐方法导致的序列偏好语义偏离, 并证实简单的詹森-香农散度过滤策略即可缓解这一问题. 在此基础上, 设计语义偏移程度感知的修复机制. REDRec 以詹森-香农散度量化增强序列的语义偏离程度, 并为不同偏离程度的样本执行基于差异化贝塔分布的嵌入表示插值操作以实现语义修复. 为了丰富增强数据的标签, REDRec 对修复后样本执行基于扰动的多轮推理, 通过自适应聚合这些推理结果生成新的软标签. 新的软标签将作为额外的数据添加到模型训练过程中. 在多个代表性骨干网络和真实世界数据集上的实验验证了 REDRec 的有效性和泛化性.

关键词: 序列推荐; 数据增强; 语义信息; 嵌入混合

中图法分类号: TP311

中文引用格式: 党翌洲, 赵楚, 姜琳颖, 马连博, 郭贵冰. 序列推荐启发式数据增强的修复和丰富. 软件学报. <http://www.jos.org.cn/1000-9825/7678.htm>

英文引用格式: Dang YZ, Zhao C, Jiang LY, Ma LB, Guo GB. Repairing and Enriching Heuristic Data Augmentation for Sequential Recommendation. Ruan Jian Xue Bao/Journal of Software (in Chinese). <http://www.jos.org.cn/1000-9825/7678.htm>

Repairing and Enriching Heuristic Data Augmentation for Sequential Recommendation

DANG Yi-Zhou, ZHAO Chu, JIANG Lin-Ying, MA Lian-Bo, GUO Gui-Bing

(Software College, Northeastern University, Shenyang 110169, China)

Abstract: Heuristic data augmentation effectively mitigates data sparsity in sequential recommendation by perturbing the original user behavior sequences and generating new ones. Compared with augmentation methods that rely on external information, complex augmentation rules, or learnable generation modules, heuristic methods perform augmentation through simple operations based on prior human knowledge, rules, or intuition, offering high efficiency and generalizability. However, existing heuristic data augmentation methods typically face two challenges. First, the augmented data often exhibits semantic bias or loss of important user behaviors compared with the original data, leading models to learn inaccurate or even erroneous user preferences. Second, existing approaches focus on augmenting sparse input behaviors, yet the model training remains constrained by the sparsity and mismatch in labeled behaviors, limiting improvements in recommendation performance. To address this issue, this study proposes a framework for repairing and enriching heuristic data augmentation in sequential recommendation (REDRec). Specifically, this study first validates the semantic deviation in sequence preferences caused by heuristic recommendation methods through empirical studies and demonstrates that a simple Jensen-Shannon

* 基金项目: 国家自然科学基金 (62576083)

收稿时间: 2026-01-14; 修改时间: 2026-02-20; 采用时间: 2026-03-23; jos 在线出版时间: 2026-07-01

divergence filtering strategy can mitigate this issue. On this basis, a repair mechanism that is sensitive to the degree of semantic shift is designed. REDRec quantifies the semantic deviation of augmented sequences using Jensen-Shannon divergence and performs an embedding representation interpolation based on differentiated beta distributions for samples with varying levels of deviation to achieve semantic repair. To enrich the labels of augmented data, REDRec performs perturbation-based multi-round inference on repaired samples and generates new soft labels by adaptively aggregating these inference results. The new soft labels are added as additional data during model training. Experiments on multiple representative backbone networks and real-world datasets validate the effectiveness and generalization of REDRec.

Key words: sequential recommendation; data augmentation; semantic information; embedding mixing

序列推荐通过从用户的历史交互序列中学习其偏好和行为模式,为用户提供个性化的下一项交互建议^[1].由于与真实推荐场景的高度一致性,序列推荐在近年来受到了高度关注与广泛研究,研究者们提出了许多基于不同架构的序列推荐模型^[2-4].训练这些序列推荐模型需要海量高质量的用户行为作为支撑.然而,与可通过众包标注或互联网文档获取的图像文本数据不同,个性化推荐依赖用户的个性化行为数据^[5].但是,用户通常只会在平台上与少量物品产生互动,加之跨平台数据互通的限制和个人隐私保护的规定^[6-8],众多历史数据难以被收集或用于序列推荐模型的训练.因此,数据稀疏性问题严重制约了序列推荐模型的性能.

为了缓解这一问题,研究者们提出了众多序列推荐数据增强方法并取得了显著的进展.数据增强是一种在计算机视觉和自然语言处理领域得到广泛应用并被证明有效的技术手段^[9,10].它通过在原始数据上施加一定扰动或是直接利用生成模块,产生额外的训练样本,有效缓解了数据稀疏性.在序列推荐领域,现有数据增强方法的研究可以主要分为两条路线:(1)基于启发式操作的数据增强通过滑动窗口^[11]、随机掩码^[12]、随机插入^[13]等简单操作,对原始序列施加扰动来生成新的序列.这些操作通常基于人类先验知识、特定规则或直觉而设计.(2)基于模型的数据增强则通过训练精心设计的数据生成模块来扩充原始数据.例如利用双向 Transformer 学习长用户序列的行为模式来延长较短的用户序列^[14];或通过去噪扩散模型直接生成新的行为序列用于模型训练^[15].相较于需要额外参数、复杂训练流程和部署过程的基于模型的增强方法,启发式增强方法不会引入额外的模型参数,不需要修改原始模型的训练流程,部署更为简单.其凭借简洁性、高效性和强大的泛化能力,已经得到了广泛的应用^[16-18].

尽管启发式增强具备众多优势,其仍然面临以下两个局限:(1)序列语义信息破坏:启发式数据增强方法的随机扰动操作可能会破坏原始序列数据的因果顺序,或直接导致重要语义信息被削弱,从而建模出错误的顺序模式或不完整的用户偏好学习.图1给出了一个例子.一位对摄影和电子产品感兴趣的白领工作者依次购买了“电脑”“西装”“照相机”“鼠标”和“镜头”等物品.在数据增强过后,重新排序操作破坏了购买序列中的逻辑顺序,例如先购买“相机”再购买“镜头”,先购买“电脑”再购买“鼠标”.同时,裁剪生成的序列削弱了序列中包含的用户对摄影产品的兴趣.(2)标签的稀缺与不匹配问题:现有的启发式数据增强方法在增强时往往只专注于对模型训练时的输入序列进行增强,而保持原有的标签不变.然而,这一做法不仅使标签数据仍然稀疏,无法为模型提供充足的监督信号,还可能导致输入序列与标签的不匹配问题.在图1的例子中,原始用户序列在模型训练阶段的标签是“耳机”,而在数据增强操作过后,新生成的序列所包含的用户偏好已经发生了轻微改变,这使再使用“耳机”作为标签行为则可能使模型学习到不准确的偏好知识和行为关联,进而带来次优的推荐性能.因此,为不同的增强序列提供丰富且合适的标签是有必要的.

本文首先通过实证研究验证已有启发式数据增强方法对用户序列语义信息的破坏.结果表明,在一些例子中,应用启发式数据增强操作后不仅不会带来推荐性能的提升,反而会损害模型的推荐准确性.通过对序列嵌入的进一步可视化分析表明,在性能欠佳的案例中,启发式数据增强方法生成的序列数据在分布上与原始数据存在显著偏差.上述结果验证了启发式数据增强操作对序列语义信息和偏好知识的破坏问题.受生成式对抗网络中生成嵌入与真实嵌入差异测量方法的启发^[19],可以将原始序列嵌入和增强序列嵌入视为语义与偏好的概率分布,进而采用詹森-香农散度 (Jensen-Shannon divergence, JSD) 量化增强序列和原始序列的偏离程度.在此基础上,本文移除增强数据中偏离程度过高的样本,使用剩余数据正常参与模型训练.此时,模型性能有所改善,特别是在先前因启发式数据增强而导致性能损害的例子.这一结果进一步说明了对启发式数据增强进行修复的潜力和重要性.

针对已有方法的问题, 同时基于实证研究中的结果和分析, 本文提出了一个称为“修复和丰富序列推荐启发式数据增强 (REDRec)”的方法, 用于解决启发式数据增强序列语义信息破坏和标签的稀缺与不匹配问题. REDRec 包含两个主要模块, 分别是修复模块和丰富模块. 对于修复模块, 其核心思路是通过嵌入混合^[20]操作在表示层面修复受损的偏好和语义信息. 具体来说, 本文利用实证研究中采用的詹森-香农散度量化增强序列的偏离程度, 根据偏离程度不同, 本文在嵌入混合时从不同的贝塔分布中采样混合权重. 换句话说, 对于偏离程度较高的增强序列, 在混合时就会赋予原始序列嵌入表示更高的权重, 偏离程度较低的序列则反之. 修复模块在保留增强序列所包含的多样变化的基础上补充了丢失的重要语义和偏好知识. 对于丰富模块, 其核心思路是利用模型在修复样本上预测得到的标签作为额外的监督信号参与模型训练. 具体来说, REDRec 对经过编码器编码的修复后嵌入执行多次独立的随机丢弃操作以生成扰动嵌入. REDRec 利用模型分别在这些扰动嵌入上执行正常的推理过程并得到预测的交互概率分布, 同时, 模型会计算对该预测结果的置信度. 之后, 根据置信度自适应地聚合这些概率分布得到最终预测结果, 即丰富后的标签. 丰富后的标签会连同经过修复后的嵌入表示一同作为新的训练样本. 最后, 所提出的方法采用了一种两阶段训练策略, 以避免混合未经充分学习和优化的物品嵌入带来的噪声样本和模型收敛困难.



图1 现有启发式数据增强方法在语义信息和标签上的局限示例

总的来说, 本文的主要贡献可以概括总结为以下 3 个关键方面.

(1) 发现并分析已有启发式数据增强方法中语义信息被破坏和标签稀缺与不匹配问题. 在代表性序列模型和启发式数据增强方法上进行的实证研究和可视化分析进一步展示了启发式操作对原始序列语义和偏好知识的破坏, 以及对其进行修复的潜力和重要性.

(2) 提出一个模型不可知的, 名为“修复和丰富序列推荐启发式数据增强 (REDRec)”的框架. 修复模块通过嵌入混合操作在表示层面修复受损的偏好和语义信息, 丰富模块则利用模型在修复样本上预测得到的标签作为额外的监督信号参与模型训练.

(3) 在多个真实世界的数据集以及多种具有代表性的序列推荐骨干网络上进行广泛的实验, 验证了所提出方法的有效性和泛化性. REDRec 可以被无缝集成到现有的序列推荐模型和启发式方法中, 在保留原始序列中重要知识的同时引入增强操作的多样变化.

本文第 1 节回顾与本文相关的研究工作, 包括序列推荐及数据增强用于序列推荐. 第 2 节介绍序列推荐和嵌入混合的符号、问题定义等预备知识. 第 3 节对已有问题进行实证研究和可视化分析. 第 4 节详细阐述本文提出的方法及其设计细节. 第 5 节展示大量实验结果, 以验证所提方法的有效性和泛化性. 最后, 第 6 节对本文的研究内容进行总结与展望.

1 相关工作

本节将从两个主要方面对本文的相关工作进行详细阐述: 一是传统的序列推荐方法, 二是基于数据增强的序

列推荐方法.

1.1 序列推荐方法

顺序推荐 (sequential recommendation, SR) 通过学习用户交互序列中物品间的转换模式和偏好知识, 预测下一个可能的交互. 早期研究利用马尔可夫链, 根据用户最近几次交互来预测后续交互. 例如, He 等人^[21]提出了一种名为 Fossil 的方法用于预测个性化序列行为, 该方法融合了基于相似性的模型与马尔可夫链. 基于马尔可夫链的方法虽然能捕捉用户的短期偏好, 但难以有效建模长时间跨度的依赖关系. 为此, 基于循环神经网络的解决方案应运而生. 循环神经网络既能捕捉偏好短期变化, 又能反映长期趋势. 例如, Hidasi 等人^[22]提出使用门控循环单元来建模用户偏好的演变过程. 在此基础上, Cui 等人^[23]开发了多视角循环模型, 有效缓解了冷启动问题. 该模型通过融合物品潜在特征、视觉特征和文本特征构建全面的物品表示. 为了加速计算, Qu 等人^[24]将邻近信息单元分割为块, 并在特定时间步执行记忆访问, 从而大幅减少记忆操作次数. 然而, 循环神经网络存在并行性差的缺陷, 这限制了其处理大规模数据时的效率. 针对这一问题, 研究者采用卷积神经网络来学习序列模式^[25]. 代表性研究有 Caser^[26]和 NextItNet^[27]. 前者通过将近期物品序列嵌入时间与潜在空间, 利用卷积滤波器学习序列模式, 后者则采用多层空洞卷积层架构, 能高效学习短程与长程物品依赖关系的高层次表示.

得益于 Transformer 架构^[28]的成功, 许多基于自注意力机制的序列推荐模型应运而生. 其中, Kang 等人^[3]提出的 SASRec 作为典型代表, 采用多头自注意力机制学习不同物品的重要性, 以此为用户提供下一项的个性化建议. 基于 SASRec, Li 等人^[29]在模型中同时建模物品绝对位置及其时间间隔, 使模型能够根据不同交互之间的时间差学习到更细粒度的偏好知识. 此外, Fan 等人^[30]将每个物品嵌入随机高斯分布, 并设计瓦瑟斯坦自注意力模块来表示物品间关系, 以此更好地捕捉物品之间的协作性和传递性. 此外, 赵容梅等人^[31]设计了一种并行的三支图卷积网络结构来学习跨域序列推荐中的域间的低阶协同过滤关系和域内的高阶序列依赖关系. 其次, 在三支图结构的基础上, 提出了一种改良的外部注意力机制, 以更低的开销挖掘物品序列依赖关系并将注意力权重在多个分支上共享. 在 Transformer 和自注意力架构之外, Zhou 等人^[4]发现用户行为日志数据通常包含噪声交互, 数字信号处理领域的滤波算法能有效缓解深度序列推荐模型中的噪声干扰. 受此启发, 他们提出了一种适用于序列推荐任务的全多层感知器模型, 该模型配备可学习滤波器, 称为 FMLPRec. 在此基础上, 卢晓凯等人^[5]提出了 GraphMLP-Mixer, 这一框架首先构造全局物品图来增强模型对序列间协同信号的建模, 然后将感知机-混合器架构与图神经网络结合, 得到图-感知机混合器模型对用户兴趣进行充分挖掘. Yue 等人^[32]通过引入矩阵对角化线性递归机制高效捕捉用户行为转换模式, 同时构建了加速序列推荐模型训练与推理的递归并行化框架, 显著降低了模型训练和推理所需的时间.

1.2 数据增强用于序列推荐

面对序列推荐领域普遍存在的用户行为序列的稀疏问题, 研究者们提出了多种数据增强方法并取得了显著的进展^[17]. 现有序列推荐中数据增强方法的研究可以被分为两条路线: 基于启发式的方法和基于模型的方法. 启发式的方法通过对原始序列施加扰动操作来生成新的序列. 这些操作通常基于人类先验知识、特定规则或直觉而设计. 早期方法包括滑动窗口^[26]和随机丢弃^[33], 即通过将原始序列划分为多个子序列或随机移除序列中的部分交互实现增强. 随着对比学习在序列推荐中的广泛应用^[18], 许多启发式增强算子被提出, 如随机裁切、随机掩码、随机重排序、随机插入和随机替换等^[12,13]. 研究者们利用这些算子来构建自监督信号用于对比学习, 以此挖掘更多的用户偏好知识. 在此基础上, 有许多研究者进一步改进了这些启发式算子. 例如 Dang 等人^[16]发现模型可以从时间间隔均匀的序列中更好地学习用户偏好和行为模式, 为此, 他们提出了 5 个时间间隔感知的数据增强算子, 以此将非均匀的序列转换为均匀序列. 钱忠胜等人^[7]则在此基础上进一步构建了超图卷积网络和 Transformer 编码器相结合的对偶视图, 从多视角捕捉会话间的隐藏高阶关系, 以丰富推荐结果的多样性. 此外, Chen 等人^[34]引入一个潜在变量作为增强操作来表示用户的意图, 并通过聚类来学习该潜在变量的分布函数, 之后并利用对比式自监督学习优化序列推荐模型. Dang 等人^[35]从数据的相关性和多样性角度出发提出了 M-Reorder 和 M-Substitute 两个算子, 以此在表示层面为模型提供兼备两种性质的样本. 此外, 在对比学习方面, Hu 等人^[36]构建了时序多兴趣预训

练框架, 通过锚点基于的对比学习和兴趣内 (间) 对比学习等任务, 优化多兴趣表示以提升推荐性能. Xiao 等人^[37]提出了采用双对比学习策略, 通过序列增强生成视图进行增强序列对比学习, 并对齐含相同目标项的异质序列开展同目标序列对比学习, 缓解噪声与数据稀疏问题. Liu 等人^[38]通过随机掩码和添加噪声生成序列增强视图, 将原始序列与增强视图作为正样本, 将同批次其他用户的增强表示作为负样本, 以对比损失优化, 提升用于序列推荐的个性化提示的表征质量.

由于启发式数据增强的随机操作可能产生低质量的增强数据以及可控性较低, 研究者们提出了可学习的复杂增强模块来提高增强数据的质量. Liu 等人^[15]采用扩散过程生成序列数据. 他们提出了一种序列式 U-Net 在预测噪声的同时捕捉序列信息. 此外, 还设计了两项约束条件以确保生成数据的质量. Wang 等人^[39]运用反事实推理对原始行为数据进行适当的替换操作以生成增强序列. Liu 等人^[40]首先训练了一个逆向 Transformer 来扩展数据集中较短的用户行为序列, 随后利用扩展后的短序列进一步提升模型预测用户未来交互的性能. 沿袭这一思路, Jiang 等人^[14]通过引入正向学习约束以此在逆向延长短序列时有效捕捉上下文信息. 由于正向约束与序列推荐任务相契合, 它能弥合逆向增强与正向推荐之间的差距. 此外, Wang 等人^[41]定义了核心用户与休闲用户 (即分别具有长序列和短序列的用户), 他们从核心用户的行为数据中学习增强策略来生成成交交互序列来模拟休闲用户的行为, 从而在不牺牲核心用户体验的前提下提升推荐模型在休闲用户上的表现. Yin 等人^[42]提出了一个以数据为中心的序列重生成框架来优化数据集中的序列信息, 提高模型的推荐准确性. 为了增强序列推荐模型应对分布偏移的能力, Zhou 等人^[43]提出分布鲁棒序列推荐框架, 通过样本重加权模拟不同分布并优化最坏情况损失以应对分布偏移, 又采用硬和软聚类为同簇样本分配统一权重降低优化难度. Liao 等人^[44]从不变性学习视角, 通过子序列提取与混合模拟多样环境, 结合对抗训练和不变风险最小化, 学习稳定用户偏好以缓解分布偏移. 类似地, Lin 等人^[45]利用多标签类别信息量化兴趣漂移, 通过生成漂移表示和解纠缠优化分布, 学习用户兴趣表示以应对兴趣漂移.

与上述研究不同, 本文从实证研究所发现的问题出发, 通过精心设计的修复和丰富模块, 改善了已有启发式数据增强方法破坏序列语义信息和标签稀缺与不匹配问题, 进一步提高了启发式数据增强方法的实际应用价值和泛化能力.

2 基础知识

本节首先在第 2.1 节介绍序列推荐的符号和问题定义, 接着在第 2.2 节介绍最常用的启发式数据增强算子.

2.1 序列推荐

假设有用户集合 $\mathcal{U}$ 和物品集合 $\mathcal{V}$. 每个用户 $u \in \mathcal{U}$ 都关联着一条按时间顺序由远及近排列的交互物品序列 $s_u = [v_1, v_2, \dots, v_{|s_u|}]$, 其中 $v_j \in \mathcal{V}$ 表示用户 u 在时间步 j 交互过的物品, $|s_u|$ 表示交互序列的长度. 给定这样的一条交互序列 s_u , 序列推荐的目标是准确预测用户 u 在时间步 $|s_u| + 1$ 最有可能交互的物品 v^* . 该过程可用公式 (1) 来表述:

$$\arg \max_{v^* \in \mathcal{V}} P(v_{|s_u|+1} = v^* | s_u) \quad (1)$$

公式 (1) 还可以被解释为计算所有候选物品的概率, 并选择最高概率者进行推荐.

2.2 启发式数据增强

启发式数据增强 (heuristic data augmentation, HDA) 方法会对原始数据序列 s_u 施加特定的扰动或变换, 从而生成新的增强序列 s_{ua} , 这种过程可以表示为 $s_{ua} = \text{HDA}(s_u)$. 本研究采用了 5 种被广泛使用且已被证明有效的启发式数据增强算子, 它们分别是: 裁剪 (Crop, C)、掩码 (Mask, M)、重排序 (Reorder, R)、替换 (Substitute, S) 以及插入 (Insert, I). 为了便于后续说明, 本文使用这些操作英文名称的第 1 个大写字母来表示它们. 在这 5 种算子中, 前 3 种算子 C、M、R 是由 Xie 等人^[12]提出的, 后来 Liu 等人^[13]又在此基础上引入了替换和插入算子, 最终形成了 C、M、R、S、I 这一算子组合. 本研究也采用了 CMR 和 CMRSI 这两种被广泛使用的组合方式^[16,46]. 给定一条原始序列 s_u , 接下来, 将分别介绍这 5 个算子的细节. 在这里, 用 α 来表示算子在增强过程中所需的强度超参数.

裁切 (C): 从原始序列 s_u 中随机选取一段长度为 $L = |s_u| \times \alpha$ 的连续子序列. 其中 $\alpha \in (0, 1)$ 是决定子序列长度的超参数^[12]. 这一过程可以用公式表示为:

$$s_u^a = \text{HDA}_C(s_u, \alpha) = [v_i, v_{i+1}, \dots, v_{i+L-1}] \quad (2)$$

重排序 (R): 从原始序列 s_u 中随机选取一段连续子序列 $[v_i, \dots, v_{i+r-1}]$ 并随机打乱, 使其变为 $[v'_i, \dots, v'_{i+r-1}]$ ^[12]. 这一过程可以用公式表示为:

$$s_u^a = \text{HDA}_R(s_u, \alpha) = [v_1, v_2, \dots, v'_i, \dots, v'_{i+r-1}, \dots, v_{|s_u|}] \quad (3)$$

其中, 子序列长度的计算方法与裁剪算子相同.

掩码 (M): 从原始序列 s_u 中随机选择 α 比例 ($0 < \alpha < 1$) 的物品进行掩码处理^[12]. 这一过程可以用公式表示为:

$$s_u^a = \text{HDA}_M(s_u, \alpha) = [v'_1, v'_2, \dots, v'_{|s_u|}] \quad (4)$$

其中, 如果 v_i 是被选中的物品, 那么 v'_i 将被替换为特殊的掩码标记; 否则 $v'_i = v_i$.

替换 (S): 将原始序列中的一部分物品替换为相似物品^[13]. 替换比例为 α ($0 < \alpha < 1$). 这一过程可以用公式表示为:

$$s_u^a = \text{HDA}_S(s_u, \alpha) = [v'_1, v'_2, \dots, v'_{|s_u|}] \quad (5)$$

其中, 当 v_i 是被选中的物品时, v'_i 会被替换为与 v_i 相关的物品; 否则 $v'_i = v_i$. 这个相关物品是根据相关性得分或物品嵌入表示的相似度来确定的. 由于相似物品的选择并不是本文的重点, 这里不再详细介绍.

插入 (I): 向原始序列 s_u 中随机插入一部分物品^[13]. 插入比例为 α ($0 < \alpha < 1$). 这一过程可以用公式表示为:

$$s_u^a = \text{HDA}_I(s_u, \alpha) = [v_1, v_2, \dots, v_{x_d}, v'_{x_d}, v_{x_d+1}, \dots, v_{y_d}, v'_{y_d}, v_{y_d+1}, \dots, v_{|s_u|}] \quad (6)$$

其中, v'_{x_d} 和 v'_{y_d} 指的是与插入位置相邻物品相似的物品. 在这里, 获取这些相似物品的方法与替换操作相同.

3 实证研究

本节将通过一个实证研究揭示已有启发式数据增强方法导致的语义信息和偏好知识被破坏的问题. 此外, 将通过实验证实基于詹森-香农散度的简单过滤策略即可缓解这一问题, 为后续利用其在 REDRec 中衡量量化增强序列和原始序列的偏离程度打下基础.

3.1 已有启发式数据增强算子的表现

首先, 本文将通过实验验证现有启发式数据增强方法对重要语义信息和偏好知识的破坏效果. 首先, 选取在第 2.2 节所述的两种启发式数据增强算子组合 CMR 和 CMRSI, 并将其应用于两个代表性序列推荐模型 GRU4Rec^[22] 和 SASRec^[3] 中. 实验数据集采用来自亚马逊购物平台的 Beauty 与 Sports 数据集^[47]. 本文采用 $H@10$ 、 $H@20$ 、 $N@10$ 和 $N@20$ 作为评估指标, 这些指标数值越大, 表明推荐准确性越高. 更多实验细节将在第 5 节介绍. 表 1 展示了实验结果.

可以观察到, 启发式数据增强方法能显著提升 GRU4Rec 模型的性能表现, 这进一步验证了启发式方法的有效性和潜力. 但是, 在 SASRec 模型上, 启发式数据增强方法带来的改进效果相当有限. 在绝大多数案例中, 添加启发式数据增强方法后模型的推荐准确性反而不如基础模型. 因此, 仅使用原始数据可能不足以使 GRU4Rec 模型得到充分训练和收敛, 因此添加启发式数据增强方法后模型性能提升显著. 然而, SASRec 模型对序列语义和偏好信息的变化更为敏感, 启发式数据增强方法破坏了原始序列的重要语义特征和偏好信息, 导致 SASRec 学习到不准确甚至是错误的用户偏好, 最终造成推荐性能下降. 此外, 从单一算子上来看, Substitute 和 Insert 整体表现较好, 因为它们在进行替换或插入操作时会考虑操作位置物品与候选物品的相似性, 在一定程度上确保了增强序列的合理性. 同时, 替换和插入操作也尽可能地保留了原始序列的序列模式, 使模型能够准确捕捉用户的偏好知识. 而对于 Crop、Reorder 和 Mask 这 3 个算子, 它们在部分案例中表现较好. 但是, Reorder 在重排序的过程中可能会打乱原本的序列模式, Crop 和 Mask 的裁切和掩码操作可能导致最能代表用户偏好的关键交互丢失. 因此, 这 3 个算子在多数案例中的表现逊于 Substitute 和 Insert.

表 1 已有启发式数据增强算子在代表性模型上的表现

模型	方法	Beauty					Sports				
		H@10	N@10	H@20	N@20	提升幅度 (%)	H@10	N@10	H@20	N@20	提升幅度 (%)
GRU4Rec	基础模型	0.0404	0.0201	0.0663	0.0266	—	0.0167	0.0080	0.0292	0.0111	—
	+Crop	0.0420	0.0213	0.0695	0.0284	5.38	0.0176	0.0084	0.0315	0.0119	6.37
	+Reorder	0.0410	0.0205	0.0680	0.0272	2.07	0.0175	0.0088	0.0319	0.0117	7.36
	+Mask	0.0422	0.0211	0.0692	0.0281	4.86	0.0180	0.0087	0.0322	0.0122	9.18
	+Substitute	0.0444	0.0231	0.0729	0.0297	11.61	0.0196	0.0092	0.0335	0.0132	16.50
	+Insert	0.0450	0.0234	0.0736	0.0301	12.99	0.0187	0.0087	0.0340	0.0129	13.35
	+CMR	0.0426	0.0218	0.0712	0.0289	7.49	0.0201	0.0099	0.0344	0.0135	20.88
	+CMRSI	0.0477	0.0241	0.0757	0.0312	17.36	0.0218	0.0108	0.0372	0.0147	31.34
SASRec	基础模型	0.0596	0.0318	0.0910	0.0396	—	0.0311	0.0171	0.0484	0.0215	—
	+Crop	0.0583	0.0315	0.0892	0.0383	-2.10	0.0315	0.0169	0.0504	0.0219	1.53
	+Reorder	0.0540	0.0276	0.0829	0.0347	-10.97	0.0309	0.0173	0.0492	0.0217	0.78
	+Mask	0.0565	0.0304	0.0868	0.0370	-5.20	0.0315	0.0173	0.0518	0.0222	3.18
	+Substitute	0.0627	0.0337	0.0966	0.0431	6.54	0.0307	0.0165	0.0470	0.0205	-3.08
	+Insert	0.0616	0.0335	0.0949	0.0416	4.51	0.0305	0.0167	0.0461	0.0213	-2.49
	+CMR	0.0560	0.0297	0.0865	0.0368	-6.16	0.0318	0.0172	0.0522	0.0224	3.72
	+CMRSI	0.0594	0.0315	0.0887	0.0383	-1.77	0.0309	0.0171	0.0468	0.0210	-1.57

注: 提升幅度是相对于基础模型的改进百分比的平均值

3.2 可视化分析

为了进一步分析启发式数据增强方法导致 SASRec 模型性能下降的原因, 本文从 Beauty 和 Sports 数据集中分别随机抽取 256 条用户序列, 采用 t-SNE 算法^[48]对原始序列及其增强序列在模型中的嵌入表示进行可视化分析. 此处使用的 SASRec 是经过充分训练的, 而不是仅经过初始化的模型. 可视化结果如图 2 所示. 可以观察到 CMR 和 CMRSI 生成的增强序列在语义分布上与原始序列存在显著差异, 这种差异在 Beauty 数据集上更为明显. 该现象与表 1 中的性能表现相吻合, 即启发式数据增强方法在 Beauty 数据集上的表现相较于在 Sports 数据集上更差. 虽然数据增强方法需要一定程度的增强数据多样性, 但显著差异可能适得其反^[35]. 综上所述, HDA 方法会削弱原始序列的重要语义信息和偏好知识, 导致模型学习到不准确甚至错误的用户偏好, 进而无法生成准确的推荐结果.

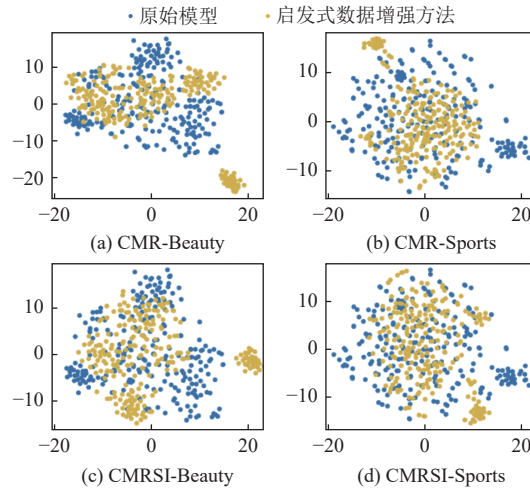


图 2 现有启发式数据增强方法在语义信息和标签上的局限示例

3.3 基于詹森-香农散度的过滤策略

受到生成式对抗网络中生成嵌入与真实嵌入差异测量方法的启发^[19],将原始序列嵌入和增强序列嵌入视为语义与偏好的概率分布,进而采用詹森-香农散度(JSD)量化增强序列和原始序列的偏离程度.具体计算过程将在第4节中介绍.在此基础上,本文移除增强数据中偏离程度排在前 Q 的样本,使用剩余增强数据正常参与训练.其中, Q 在{5%, 10%, 15%, 20%, 25%}中调优.结果展示在表2中.

表2 已有启发式数据增强算子添加基于詹森-香农散度过滤前后的性能表现

模型	方法	Beauty					Sports				
		H@10	N@10	H@20	N@20	提升幅度(%)	H@10	N@10	H@20	N@20	提升幅度(%)
GRU4Rec	CMR	0.0426	0.0218	0.0712	0.0289	—	0.0201	0.0099	0.0344	0.0135	—
	+JSD	0.0439	0.0227	0.0739	0.0301	3.78	0.0207	0.0103	0.0368	0.0139	4.24
	CMRSI	0.0477	0.0241	0.0757	0.0312	—	0.0218	0.0108	0.0372	0.0147	—
	+JSD	0.0493	0.0252	0.0780	0.0324	3.70	0.0222	0.0111	0.0381	0.0150	2.27
SASRec	CMR	0.0560	0.0297	0.0865	0.0368	—	0.0318	0.0172	0.0522	0.0224	—
	+JSD	0.0621	0.0326	0.0929	0.0411	9.94	0.0332	0.0181	0.0540	0.0231	4.05
	CMRSI	0.0594	0.0315	0.0887	0.0383	—	0.0309	0.0171	0.0468	0.0210	—
	+JSD	0.0629	0.0335	0.0947	0.0417	6.97	0.0320	0.0175	0.0526	0.0221	5.88

可以观察到,在利用詹森-香农散度过滤偏离的样本后,所有启发式数据增强方法的性能都有所提升,最高可以接近10%.结合表1和表2,在添加过滤后,启发式数据增强方法导致模型性能下降的情况被明显缓解.这一结果表明利用詹森-香农散度衡量增强序列和原始序列在语义信息与偏好知识上偏离程度的有效性.此外,当过滤掉过于偏离的序列时,模型能够从剩余序列中进一步学习用户偏好,获得更好的推荐性能.但是,简单的过滤可能无法有效滤除所有偏离的序列,而一些原本没有被过滤掉的序列可能也会包含偏离的信息.因此,提出一种能够有效修复增强序列中受到损害的语义的方法是必要且具有价值的.

4 序列推荐启发式数据增强的修复和丰富方法 REDRec

基于实证研究中的结果和分析,在本节中,将介绍所提出的 REDRec 框架.如图3所示,该框架包含两个增强模块:修复模块和丰富模块.随后,将详细说明 REDRec 的训练流程.最后,将对 REDRec 与现有方法进行讨论和对比分析.

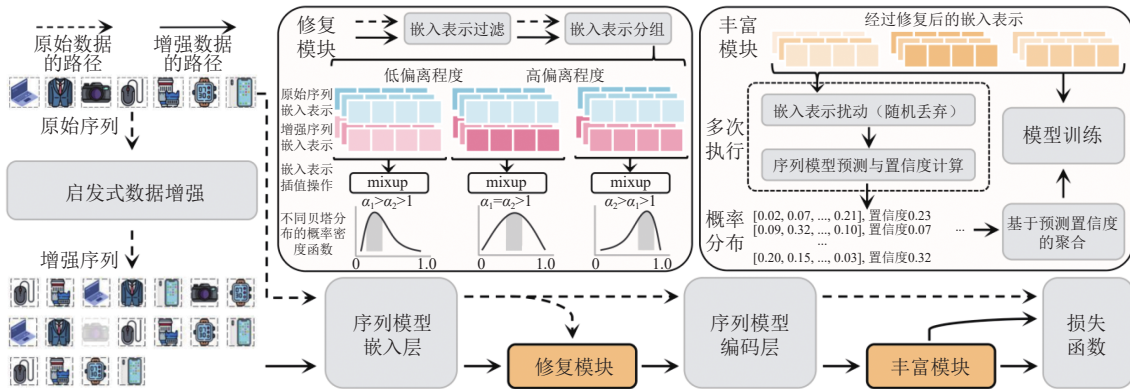


图3 提出的 REDRec 的整体框架示意图

4.1 修复模块

为修复启发式数据增强方法在对序列施加扰动过程中导致其受损的语义信息和偏好知识,丰富模块在物品序列嵌入表示层面对原始序列与增强序列进行混合,生成混合样本.混合样本既包含扰动引入的多样变化,又保留原

始语义信息和偏好知识. 在标准序列推荐的范式中, 通常维护一个物品嵌入矩阵 $\mathbf{M}_V \in \mathbb{R}^{|\mathcal{V}| \times D}$, 该矩阵将高维独热物品索引投影为低维密集嵌入. 给定原始用户序列 $s_u = [v_1, v_2, \dots, v_{|s_u|}]$, 通过在 $\mathbf{M}_V$ 上执行的查找操作 (look up) 可获得物品的嵌入表示序列 $E_u = [e_1, e_2, \dots, e_{|s_u|}]$. 这里需要注意的是, 本文在叙述中省略了对短序列的填充处理以及对长序列省略截断处理. 给定原始序列 s_u 及其通过启发式数据增强方法生成的增强序列 s_{ua} , 可分别获得其序列嵌入表示 E_u 和 E_{ua} .

在获得了两种序列的嵌入表示后, 本文设计了一种偏离程度感知的嵌入表示混合策略. 直观来说, 本文希望对于偏离程度高的增强序列, 在混合时原始序列的嵌入表示能占更高的权重, 而偏离程度低的序列则反之. 遵循实证研究中的思路, 本文首先使用詹森-香农散度衡量增强序列和原始序列在语义信息与偏好知识上偏离程度, 这一过程可以用公式表示为:

$$Dev = \frac{1}{2} \left(E_u \log \frac{2E_u}{E_u + E_{ua}} + E_{ua} \log \frac{2E_{ua}}{E_u + E_{ua}} \right) \quad (7)$$

在得到 Dev 值后, 将其除以 $\log 2$ 以映射至 $[0, 1]$. Dev 越接近 1, 表示增强序列较原始序列偏离程度越大. 遵循实证研究中探索的结果, 对于一个批次的增强数据, 本文首先过滤掉偏离程度前 Q 的增强序列. 之后, 为不同偏离程度的剩余增强序列提供不同的混合权重, 将偏离值从 0 到 1 平均分为 T 个区间, 每个区间将对应一个专门的贝塔分布, 用于在混合时采样混合权重. 接下来, 将介绍如何生成这些贝塔分布.

传统的样本混合操作 (mixture) 最初在视觉领域被提出作为一种简单有效的数据增强方法^[20]. 该方法通过线性插值两个原始样本生成新样本. 例如当原始样本为 x_i 、 x_j 时, 插值过程可以表示为 $\tilde{x} = \lambda \cdot x_i + (1 - \lambda) \cdot x_j$. 其中, $\tilde{x}$ 是插值样本, 而插值权重 λ 从一个贝塔分布中采样得到, 即 $\lambda \sim \text{Beta}(\alpha_1, \alpha_2)$. 在通常情况下, $\alpha_1 = \alpha_2 < 1$. 然而, 当 α 值较小时 (通常为 0.2 或 0.3) 时, 贝塔分布会在 $\lambda = 0$ 和 $\lambda = 1$ 处出现峰值, 这意味着混合后的嵌入表示要么完全由原始序列的表示主导, 要么完全由增强序列的表示主导. 前者无法引入足够的变化, 而后者则只能产生微弱的修复效果. 因此, 为了避免上述情况发生, 同时为不同偏离程度的增强序列提供不同的混合权重分布, 将初始化多个 α_1 和 α_2 均大于 1 的贝塔分布.

在图 3 修复模块的下半部分展示了当 $\alpha_1 = \alpha_2 > 1$ 时, 贝塔分布的概率密度函数图像. 此时混合后的嵌入中的原始序列与增强序列具有近似相等的权重. 当 α_1 与 α_2 大小不同时, 权重分布会向一侧偏移, 且差异越大, 偏移幅度越大. 因此, 只需调整 α_1 值即可改变混合权重 λ 的分布, 从而为不同偏离程度的增强序列提供不同的混合权重. 具体来说, 给定超参数 $\alpha_a > \alpha_b > 1$ (α_a 和 α_b 为 α_1 和 α_2 的初始值), 由区间数量 T 可以计算得到一个步长 S , 其中 $S = (\alpha_a - \alpha_b) / T$. 根据步长 S 这些 α 值对可以具体表示为:

$$\{(\alpha_a, \alpha_b), (\alpha_a + S, \alpha_b - S), (\alpha_a + 2S, \alpha_b - 2S), \dots, (\alpha_a + TS, \alpha_b - TS)\} \quad (8)$$

根据这些 α 值对, 可以得到 T 个贝塔分布, 从左到右依次对应偏离值从 0 到 1 的 T 个区间. 在得到了混合所需的贝塔分布后, 表示混合过程将按公式 (9) 执行:

$$E_u^R = \lambda \cdot E_u + (1 - \lambda) \cdot E_{ua} \quad (9)$$

其中, E_u^R 是修复后用于后续训练的代表.

上述过程实现了为偏离值不同的增强序列分配不同的贝塔分布以采样混合权重. 偏离程度越高的序列 (Dev 值越大), 将会被分配到 α_1 越大的贝塔分布. 如图 3 修复模块的下半部分所示, α_1 越大, 则混合权重 λ 的概率密度函数会越向 1, 此时, 通过公式 (9) 进行嵌入混合时, 原始序列的占比就会更高, 进而更大程度地修复增强序列中损失的语义信息和偏好知识. 偏离程度越低的序列 (Dev 值越小), 将会被分配到 α_2 越大的贝塔分布. α_2 越大, 则混合权重 λ 的概率密度函数会越向 0, 此时, 通过公式 (9) 进行嵌入混合时, 原始序列的占比就会更小, 进而保留增强序列中引入的多样的变化, 进一步提高模型的泛化能力.

4.2 丰富模块

丰富模块的核心目标是解决现有启发式数据增强中标签稀缺与输入-标签不匹配的问题. 该模块以修复模块输出的嵌入 E_u^R 表示为基础, 通过多轮扰动推理与置信度感知的聚合, 生成与增强序列语义匹配的软标签, 为模型

训练提供额外且高质量的监督信号. 在得到修复后的嵌入 E_u^R 后, 将其输入序列行为编码器并获得输出表示 H_u^R . 为充分挖掘修复后嵌入中的偏好信息, 同时模拟预测场景的多样性, 本文对 H_u^R 执行多次独立的随机丢弃 (*Dropout*) 操作, 生成 N 个扰动表示. 随机丢弃通过随机屏蔽表示中的部分值, 能够模拟不同的特征选择视角, 迫使模型关注序列中更核心的偏好信号, 而非依赖局部冗余信息, 这一过程可以用公式 (10) 表示:

$$E_u^{R(n)} = \text{Dropout}(E_u^R), n = 1, 2, \dots, N \quad (10)$$

将 N 个扰动表示分别输入序列推荐模型的预测层, 执行正常的推荐推理过程, 得到每个扰动嵌入对应的物品交互 $P^{(n)} = [p_1^{(n)}, p_2^{(n)}, \dots, p_v^{(n)}, \dots, p_{|V|}^{(n)}]$, 其中 $p_v^{(n)}$ 表示用户在第 n 次扰动视角下与物品 v 交互的预测概率. 为量化每个预测结果的可靠性, 本文引入了一个预测置信度指标 $\text{Conf}(P^{(n)})$, 基于概率分布的熵值计算:

$$\text{Conf}(P^{(n)}) = 1 - \frac{1}{\ln|V|} \sum_{v \in V} p_v^{(n)} \ln p_v^{(n)} \quad (11)$$

其中, 熵值越小表示概率分布越集中, 预测结果越明确, 置信度越高 (取值范围为 $[0, 1]$). 这一设计能够在聚合时降低不可靠的预测结果的权重, 避免低质量标签引入噪声.

为生成最终的软标签, 首先基于各预测结果的置信度进行加权聚合, 使高置信度的预测结果占据更大权重, 确保软标签的可靠性. 本文首先对 N 个置信度进行归一化处理, 得到权重系数 $w^{(n)}$:

$$w^{(k)} = \frac{\exp(\text{Conf}(P^{(k)}))}{\sum_{i=1}^K \exp(\text{Conf}(P^{(i)}))} \quad (12)$$

在这里, 本文采用指数函数放大置信度差异, 旨在强化高可靠预测对聚合后结果的贡献. 随后, 通过加权求和得到最终的软标签 P^{soft} :

$$P^{\text{soft}} = \sum_{k=1}^K w^{(k)} \cdot P^{(k)} \quad (13)$$

该软标签是修复后表示的多个可靠视角预测的融合结果, 既保留了不同视角下的偏好信息, 又通过置信度加权过滤了噪声, 能够准确地反映修复后序列对应的用户偏好. 生成的软标签 P^{soft} 与修复后的嵌入 H_u^R 共同构成新的训练样本. 在这里, 软标签作为对原始标签的补充, 提供了完整的物品交互概率分布, 包含更丰富的偏好监督信息, 缓解了标签稀疏性. 此外, 它基于修复后的嵌入动态生成, 与输入序列的语义高度匹配.

4.3 模型训练

在本节, 将详细介绍 REDRec 的训练流程. 基于先前研究^[3,16,32], REDRec 采用常用的二元交叉熵 (binary cross entropy, *BCE*) 损失函数用于优化序列推荐模型. 这一过程可以用公式 (14) 来表示:

$$\mathcal{L}_{\text{main}} = \text{BCE}(H_u, E_u^+, E_u^-) = -[\log(\sigma(H_u \cdot E_u^+)) + \log(1 - \sigma(H_u \cdot E_u^-))] \quad (14)$$

其中, H_u 、 E_u^+ 和 E_u^- 分别表示经过编码后的原始序列表示、正样本表示和负样本表示. $\sigma(\cdot)$ 为 Sigmoid 函数. 对于修复模块, 损失计算基于 H_u^R 和原始正负样本表示:

$$\mathcal{L}_{\text{repair}} = \text{BCE}(H_u^R, E_u^+, E_u^-) \quad (15)$$

对于丰富模块, 为了充分利用预测得到的软标签, REDRec 采用了交叉熵 (cross entropy, *CE*) 损失:

$$\mathcal{L}_{\text{enrich}} = \text{CE}(H_u^R, P^{\text{soft}}) = - \sum_{v=1}^{|V|} p_v^{\text{soft}} \cdot \log \left(\frac{\exp(H_u^R \cdot e_v)}{\sum_{i=1}^{|V|} \exp(H_u^R \cdot e_i)} \right) \quad (16)$$

此外, 为了避免在训练的初始阶段混合仅经过初始化的物品嵌入所引起的噪声样本和模型收敛问题, REDRec 引入了一种两阶段的模型训练策略. 在第 1 阶段, 遵循标准的序列推荐训练流程, 主要目标是利用原始数据学习高质量物品嵌入. 在这一阶段, 仅使用公式 (14) 作为损失函数. 在第 2 阶段, 启用修复模块和丰富模块来进一步提高模型的推荐准确性:

$$L = \mathcal{L}_{\text{main}} + \mathcal{L}_{\text{repair}} + \mathcal{L}_{\text{enrich}} \quad (17)$$

需要注意的是, REDRec 没有为 $\mathcal{L}_{\text{repair}}$ 和 $\mathcal{L}_{\text{enrich}}$ 这两个学习目标设置单独的权重, 因为增强数据通常与原始数据被同等对待^[15,35]. 在第 2 阶段同时使用 3 个学习目标也是为了更充分地利用来自两个模块的所有增强样本, 以期达到更好的推荐性能.

4.4 所提方法的优势

与现有方法相比, 本文所提出的 REDRec 具有以下优势.

(1) 无可学习参数, 即插即用: REDRec 可作为增强插件无缝集成到现有启发式数据增强方法及各类序列推荐骨干网络中, 且无需引入额外的可学习参数. 相比之下, 部分基于模型的增强方法受限于骨干网络类型, 且可能引入额外学习参数. 例如, 第 2.2 节中所介绍的 Liu 等人^[15]提出的基于扩散模型的数据增强方法需要引入额外的可学习的增强模块. Xiao 等人^[37]和 Jiang 等人^[14]提出的利用双向 Transformer 延长短序列的增强方法, 仅能在同样基于 Transformer 的序列推荐模型 (如 SASRec) 上使用.

(2) 低计算开销: 得益于嵌入混合操作在数值相加时的线性复杂度, REDRec 仅产生少量计算开销, 不需要额外花费时间训练扩散模型^[15]、双向 Transformer^[14,37]等专门的增强模块. 此外, REDRec 仅需要最基础的物品索引信息即可执行增强过程, 无需额外的辅助信息.

(3) 嵌入表示层面的多样化样本: 当前绝大多数方法只能在高维离散的交互行为层面生成新的序列数据^[12-14,37], 而本文所提出的 REDRec 可在表示层面修复受损的语义信息和偏好知识, 同时丰富增强样本的标签信息, 生成多样化且高质量的增强样本, 进一步提高模型的推荐准确性.

5 实验分析

5.1 实验数据

实验采用 4 个被广泛使用的真实世界的公共数据集: Beauty、Sports、Yelp 和 Home. 这 4 个数据集覆盖了各种用户需求和偏好, 例如美容、运动、生活服务和家居等. 它们均为近几年被大量学者使用和研究的经典数据集. 遵循先前的工作^[3,16,40], 实验使用 Yelp 数据集中时间在 2019 年 1 月 1 日之后的记录. 同时, 对于所有数据集, 交互次数少于 5 次的用户和物品将被过滤. 数据集的详细统计数据汇总于表 3.

表 3 实验数所用数据集的统计信息

数据集	用户数量	物品数量	交互数量	平均序列长度	稀疏度 (%)
Beauty	22 363	12 101	198 502	8.9	99.93
Sports	35 598	18 357	296 337	8.3	99.95
Yelp	30 431	20 033	316 354	10.4	99.95
Home	66 519	28 237	551 682	8.3	99.97

(1) Beauty (<https://cseweb.ucsd.edu/~jmcauley/datasets/amazon/links.html>): 来自亚马逊购物平台的化妆品购买记录. 收集了用户在购买后对于美妆产品的评价.

(2) Sports (<https://cseweb.ucsd.edu/~jmcauley/datasets/amazon/links.html>): 来自亚马逊购物平台的运动及户外产品购买记录. 收集了用户在购买后对于该类产品的评价.

(3) Yelp (<https://www.yelp.com/dataset>): 来自生活服务及点评平台 Yelp. 该数据集包含了用户在访问生活服务地点和餐厅等位置后对其的评价.

(4) Home (<https://cseweb.ucsd.edu/~jmcauley/datasets/amazon/links.html>): 来自亚马逊购物平台的居家用品购买记录. 收集了用户在购买后对于居家用品的评价.

5.2 基线模型

实验选取 5 种具有代表性的通用序列推荐骨干模型来验证 REDRec 的有效性和泛化能力. 这些模型基于多种不同的基础架构, 包括循环神经网络、卷积神经网络、注意力机制和多层感知机.

- (1) GRU4Rec^[22]: 这是首个将门控循环单元应用于建模用户行为序列以实现顺序推荐的模型.
- (2) NextItNet^[27]: 该模型由多层空洞卷积层堆叠而成, 能够高效地从短程和长程物品依赖关系中学习表示.
- (3) SASRec^[3]: 这是一个基于 Transformer 的经典序列推荐基线模型, 它利用多头自注意力机制来完成下一个物品的推荐任务.
- (4) FMLPRec^[4]: 该模型采用全多层感知机的架构, 利用可学习的滤波模块来消除用户行为中的噪声.
- (5) LRURec^[32]: 该模型提出了基于矩阵对角化的线性递归模型, 能够有效捕捉用户行为的迁移模式.
- 此外, 本文还将 REDRec 与多个代表性的序列推荐数据增强方法进行了比较.
- (1) ASRep^[37]: 该模型采用逆向预训练 Transformer 为短用户序列生成伪先验项, 再通过微调预训练 Transformer 正向预测后续交互.
- (2) ContrastVAE^[49]: 该模型通过基于模型的方法和变分增强技术生成额外视图用于自监督对比学习, 进一步挖掘用户偏好知识.
- (3) CL4SRec^[12]: 该模型利用启发式数据增强算子生成自监督信号用于对比学习, 缓解数据稀疏性并提高模型的推荐准确性.
- (4) DiffuASR^[15]: 该模型采用扩散过程生成交互序列, 并额外添加两个约束条件以确保生成数据的质量.
- (5) DR4SR^[39]: 这是一种数据层面的框架, 旨在将原始序列优化为包含更多信息量且泛化性更高的形式.
- (6) BARec^[14]: 该模型提出双向时间数据增强方法, 结合预训练策略与知识增强的微调, 合成保留用户偏好的伪先验项, 同时深入挖掘物品间的语义关联.
- (7) BASRec^[35]: 该方法从数据的相关性和多样性角度出发提出了 M-Reorder 和 M-Substitute 两个算子, 以此在表示层面为模型提供兼备两种性质的增强数据

5.3 评估策略

实验采用留一法分割原始数据集并获取训练集、验证集和测试集, 即每条序列的最后一个物品用于测试, 倒数第 2 个物品用于验证, 剩余物品用于训练. 在预测阶段, 通过预测用户对所有在训练集中从未出现过的物品的偏好来对物品进行排序, 并将排名最靠前的 K 个 (Top- K) 物品推荐给用户. 已有文献已经采样排序会导致评估出现偏差^[50]. 评估指标包括命中率 (hit rate@ K , $H@K$) 和归一化折扣累积增益 (normalized discounted cumulative gain@ K , $NDCG@K$, 后文简称 $N@K$). 在实验结果中, 汇报 K 为 10 和 20 时两个指标的值.

- (1) 命中率 ($H@K$): 命中率直接衡量用户感兴趣的项目是否被成功推荐. 其计算公式为:

$$H@K = \frac{\sum_{u \in U} H@K(u)}{\sum_{u \in U} GT(u)} \quad (18)$$

其中, 分子表示所有用户 Top- K 推荐列表中属于测试物品的个数总和, 分母表示所有用户的真实兴趣物品总数总和.

- (2) 归一化折扣累积增益 ($NDCG@K$): 归一化折扣累积增益综合考虑了推荐列表中项目的排序质量和用户兴趣度, 通过折损方式强调高排名的相关项目, 更能体现模型对于用户偏好学习的准确性. 换句话说, 它通过将较高的分数分配给较靠前的位置来衡量所推荐物品的排列顺序是否正确. 其计算公式为:

$$DCG@K = \sum_{i=1}^K \frac{2^i - 1}{\log_2(i + 1)} \quad (19)$$

$$IDCG@K = \sum_{i=1}^K \frac{1}{\log_2(i + 1)} \quad (20)$$

$$NDCG_u@K = \frac{DCG_u@K}{IDCG_u@K} \quad (21)$$

$$N@K = \frac{\sum_{u \in U} NDCG_u@K}{|U|} \quad (22)$$

其中, t_i 代表推荐列表中物品 i 位置的相关性, 当 i 所在位置与真实值相同时为 1 否则为 0.

5.4 实现细节

实验采用原始论文中所提供的代码实现所有基线模型. 将训练批次大小、最大序列长度和物品嵌入维度分别设为 256、50 和 64. 根据原始论文中所汇报的数值和范围, 本文仔细调整了所有超参数以确保公平比较. 所有方法使用 Adam^[51] 优化器, 学习率统一设置为 0.001. 对于本文所提出的 REDRec 模型, 分别在 $\{10, 12, 14, 16\}$ 、 $\{2, 4, 6, 8\}$ 、 $\{2, 4, 6, 8, 10\}$ 、 $\{5, 10, 15, 20, 25\}$ 范围内调整用于初始化贝塔分布的 α_a 、用于初始化贝塔分布的 α_b 、区间数量 T 和扰动样本数量 N 的值. 过滤百分比 Q 的调整范围和第 3.3 节一致. 为了能够确保多个贝塔分布的差异性, 只选择 α_a 和 α_b 的大小差异大于等于 8 的组合. 选择所有实验均在配备英特尔酷睿 i7-12700 处理器和 32 GB 内存的 NVIDIA RTX 3090Ti GPU 上完成. 所有模型均基于 PyTorch 框架 1.12.1+cu116 版本实现.

5.5 整体性能提升情况

表 4 和表 5 展示了在多个具有代表性的骨干序列推荐模型上添加了启发式数据增强算子, 并在此基础上继续添加本文所提出的 REDRec 后的推荐性能情况. 从整体性能上来看, SASRec 在所有骨干模型中表现最优, 这进一步证明了注意力机制在建模序列方面的强大能力. FMLPRec 和 LRURec 也取得了具有竞争力的表现, 证明了多层感知机和线性递归单元在学习用户序列行为上的潜力. 对于启发式数据增强方法, 可以观察到它们在很多案例上带来了一定的性能提升, 但同时, 在很多案例中, 引入它们仅能带来微乎其微的推荐性能提高, 甚至损害推荐性能. 这一现象与第 3 节实证研究中所得出结论类似, 即现有启发式数据增强方法在增强过程中可能损害了序列中的重要语义信息和偏好知识, 进而导致模型学习到了不正确甚至是错误的用户偏好. 所提出的 REDRec 显著提高了所有启发式数据增强方法的性能, 在 5 个骨干模型上带来的平均提升分别为 62.87%、21.27%、39.25%、19.16% 和 33.06%. 此外, 原始的 CMRSI 算子组合往往优于原始的 CMR, 这可能是因为前者额外引入了替换 (S) 和插入 (I) 这两个更复杂的操作, 能够生成更高质量的增强数据. 在引入 REDRec 后, CMR 和 CMRSI 的性能趋于相当, 且在某些情况下 CMR 甚至能超越 CMRSI. 这表明修复和丰富模块能显著提升原始简单启发式增强操作的性能. REDRec 展现出了卓越的有效性和泛化能力.

表 4 Beauty 和 Sports 数据集上原始模型、添加启发式数据增强方法和进一步添加 REDRec 后的性能对比

方法	Beauty				Sports			
	$H@10$	$N@10$	$H@20$	$N@20$	$H@10$	$N@10$	$H@20$	$N@20$
GRU4Rec	0.0404	0.0201	0.0663	0.0266	0.0167	0.0080	0.0292	0.0111
+CMR	0.0426	0.0218	0.0712	0.0289	0.0201	0.0099	0.0344	0.0135
w/ REDRec-CMR	0.0550	0.0288	0.0844	0.0362	0.0344	0.0181	0.0528	0.0227
+CMRSI	0.0477	0.0241	0.0757	0.0312	0.0218	0.0108	0.0372	0.0147
w/ REDRec-CMRSI	0.0616	0.0329	0.0918	0.0404	0.0364	0.0196	0.0541	0.0241
NextItNet	0.0326	0.0163	0.0511	0.0210	0.0154	0.0077	0.0272	0.0106
+CMR	0.0333	0.0159	0.0533	0.0209	0.0202	0.0098	0.0346	0.0134
w/ REDRec-CMR	0.0406	0.0188	0.0614	0.0240	0.0258	0.0121	0.0429	0.0165
+CMRSI	0.0308	0.0168	0.0500	0.0210	0.0203	0.0094	0.0337	0.0132
w/ REDRec-CMRSI	0.0393	0.0183	0.0595	0.0244	0.0264	0.0119	0.0436	0.0161
SASRec	0.0596	0.0318	0.0910	0.0396	0.0311	0.0171	0.0484	0.0215
+CMR	0.0560	0.0297	0.0865	0.0368	0.0318	0.0172	0.0522	0.0224
w/ REDRec-CMR	0.0778	0.0441	0.1135	0.0531	0.0456	0.0252	0.0686	0.0310
+CMRSI	0.0594	0.0315	0.0887	0.0383	0.0309	0.0171	0.0468	0.0210
w/ REDRec-CMRSI	0.0792	0.0452	0.1135	0.0539	0.0462	0.0257	0.0670	0.0309
FMLP-Rec	0.0511	0.0290	0.0753	0.0351	0.0230	0.0123	0.0355	0.0154
+CMR	0.0612	0.0319	0.0924	0.0397	0.0336	0.0176	0.0524	0.0224
w/ REDRec-CMR	0.0716	0.0403	0.1036	0.0488	0.0382	0.0210	0.0600	0.0252
+CMRSI	0.0647	0.0329	0.0944	0.0418	0.0334	0.0188	0.0524	0.0226
w/ REDRec-CMRSI	0.0761	0.0416	0.1106	0.0506	0.0374	0.0205	0.0584	0.0248

表 4 Beauty 和 Sports 数据集上原始模型、添加启发式数据增强方法和进一步添加 REDRec 后的性能对比 (续)

方法	Beauty				Sports			
	H@10	N@10	H@20	N@20	H@10	N@10	H@20	N@20
LRURec	0.0582	0.0305	0.0922	0.0385	0.0323	0.0180	0.0506	0.0219
+CMR	0.0593	0.0311	0.0945	0.0403	0.0311	0.0176	0.0483	0.0210
w/ REDRec-CMR	0.0739	0.0381	0.1114	0.0465	0.0457	0.0251	0.0694	0.0315
+CMRSI	0.0605	0.0320	0.0939	0.0407	0.0319	0.0185	0.0519	0.0227
w/ REDRec-CMRSI	0.0763	0.0396	0.1102	0.0459	0.0470	0.0259	0.0720	0.0322

注: w/ REDRec-CMR和w/ REDRec-CMRSI表示在相应启发式数据增强方法上添加REDRec. 加粗代表REDRec在启发式数据增强方法的基础上带来了性能提升

表 5 Yelp 和 Home 数据集上原始模型、添加启发式数据增强方法和进一步添加 REDRec 后的性能对比

方法	Yelp				Home			
	H@10	N@10	H@20	N@20	H@10	N@10	H@20	N@20
GRU4Rec	0.0184	0.0086	0.0343	0.0126	0.0072	0.0034	0.0126	0.0048
+CMR	0.0233	0.0110	0.0432	0.0160	0.0086	0.0042	0.0154	0.0059
w/ REDRec-CMR	0.0352	0.0181	0.0591	0.0242	0.0190	0.0094	0.0308	0.0123
+CMRSI	0.0248	0.0115	0.0442	0.0164	0.0094	0.0046	0.0164	0.0063
w/ REDRec-CMRSI	0.0351	0.0177	0.0594	0.0238	0.0195	0.0100	0.0318	0.0130
NextIttNet	0.0222	0.0109	0.0406	0.0155	0.0070	0.0033	0.0125	0.0046
+CMR	0.0224	0.0108	0.0402	0.0152	0.0099	0.0050	0.0175	0.0069
w/ REDRec-CMR	0.0258	0.0126	0.0482	0.0180	0.0127	0.0063	0.0212	0.0086
+CMRSI	0.0238	0.0112	0.0413	0.0155	0.0083	0.0039	0.0151	0.0058
w/ REDRec-CMRSI	0.0280	0.0134	0.0492	0.0183	0.0105	0.0048	0.0180	0.0069
SASRec	0.0247	0.0120	0.0429	0.0166	0.0145	0.0076	0.0232	0.0098
+CMR	0.0240	0.0117	0.0417	0.0164	0.0179	0.0093	0.0295	0.0122
w/ REDRec-CMR	0.0337	0.0169	0.0554	0.0224	0.0239	0.0128	0.0356	0.0157
+CMRSI	0.0254	0.0127	0.0428	0.0172	0.0184	0.0096	0.0283	0.0120
w/ REDRec-CMRSI	0.0373	0.0192	0.0593	0.0244	0.0247	0.0134	0.0375	0.0166
FMLP-Rec	0.0187	0.0092	0.0337	0.0130	0.0141	0.0075	0.0206	0.0091
+CMR	0.0337	0.0171	0.0564	0.0227	0.0194	0.0104	0.0302	0.0131
w/ REDRec-CMR	0.0396	0.0202	0.0641	0.0263	0.0242	0.0133	0.0369	0.0168
+CMRSI	0.0317	0.0153	0.0507	0.0207	0.0184	0.0102	0.0266	0.0127
w/ REDRec-CMRSI	0.0361	0.0186	0.0585	0.0247	0.0234	0.0128	0.0344	0.0165
LRURec	0.0230	0.0114	0.0405	0.0159	0.0151	0.0079	0.0244	0.0102
+CMR	0.0252	0.0120	0.0436	0.0167	0.0154	0.0078	0.0252	0.0100
w/ REDRec-CMR	0.0297	0.0148	0.0536	0.0211	0.0226	0.0120	0.0336	0.0142
+CMRSI	0.0260	0.0123	0.0450	0.0171	0.0155	0.0083	0.0262	0.0107
w/ REDRec-CMRSI	0.0322	0.0155	0.0562	0.0220	0.0237	0.0125	0.0350	0.0149

注: w/ REDRec-CMR和w/ REDRec-CMRSI表示在相应启发式数据增强方法上添加REDRec. 加粗代表REDRec方法在启发式数据增强方法的基础上带来了性能提升

表 6 展示了所提出的 REDRec 与不同数据增强方法的性能比较. 表中还给出了每个基线方法的出版会议或期刊, 以及不同方法在 4 个数据集上单次迭代的总耗时. 考虑到不同数据增强方法对骨干模型的类型存在限制, 本文选择了适用所有方法的 SASRec 作为骨干模型. 可以观察到, REDRec 在几乎所有案例中取得了最优的推荐性能, 进一步证明了其有效性和竞争力. 在基线模型中 CL4SRec 利用对比学习从数据增强方法生成的视图中进一步挖掘用户的偏好知识, 取得了较好的性能. BSARec 则将增强数据的相关性和多样性作为额外约束添加到了增强过程中, 在多数案例中取得了次优的性能. 对于 REDRec, 它通过嵌入表示层面的语义信息和偏好知识修复模块, 以及基于多轮预测和聚合标签丰富模块, 进一步释放了启发式数据增强方法的潜力, 使模型从高质量的增强数据中全面且准确地学习用户的偏好, 提高推荐准确性. 从训练效率上来看, REDRec 相较其他数据增强方法并没有带来

较多额外的迭代时间开销. ContrastVAE、CLS4Rec、BARec 和 BASRec 等利用对比学习和复杂增强模块的方法虽能提高原始模型的性能, 但却会带来显著的额外训练时间开销. 因此, REDRec 在模型推荐准确性和效率上取得了双赢, 具备较强的实用价值.

表 6 不同数据增强方法与 REDRec 的性能对比

方法	出版信息	Beauty		Sports		Yelp		Home		单次迭代 总耗时 (s)
		H@10	N@10	H@10	N@10	H@10	N@10	H@10	N@10	
基础模型	ICDM 2018	0.0596	0.0318	0.0311	0.0171	0.0247	0.0120	0.0145	0.0076	16.8
ASReP	SIGIR 2021	0.0651	0.0349	0.0363	0.0195	0.0285	0.0135	0.0193	0.0097	30.5
ContrastVAE	CIKM 2022	0.0633	0.0329	0.0379	0.0211	0.0325	0.0166	0.0189	0.0105	40.6
CLS4Rec	ICDE 2022	0.0664	0.0354	0.0392	0.0209	0.0382	<u>0.0187</u>	0.0206	0.0109	35.9
DiffuASR	CIKM 2023	0.0667	0.0365	0.0356	0.0189	0.0303	0.0152	0.0182	0.0102	26.4
DR4SR	KDD 2024	0.0652	0.0349	0.0372	0.0197	0.0286	0.0142	0.0168	0.0084	29.2
BARec	TKDE 2025	0.0729	0.0397	0.0386	0.0216	0.0347	0.0171	<u>0.0231</u>	<u>0.0119</u>	37.7
BASRec	AAAI 2025	<u>0.0750</u>	<u>0.0403</u>	<u>0.0411</u>	<u>0.0231</u>	0.0311	0.0159	0.0197	0.0106	46.1
REDRec	—	0.0792	0.0452	0.0462	0.0257	<u>0.0373</u>	0.0192	0.0247	0.0134	30.3

注: 基础模型是SASRec, 加粗表示最优性能, 下划线表示次优性能

5.6 消融实验

本文进行了消融实验来验证 REDRec 中各个组件的有效性. 消融实验对比了以下 7 种变体: (1) 加权融合损失: 为 $\mathcal{L}_{\text{repair}}$ 和 $\mathcal{L}_{\text{enrich}}$ 赋予额外的权重, 总和为 1. (2) 硬标签主导的互补: 仅为 $\mathcal{L}_{\text{enrich}}$ 赋予额外的权重, 范围在 $[0, 1]$, 步长为 0.1. (3) 移除修复模块: 移除 REDRec 中的语义信息和偏好知识的修复模块. (4) 固定贝塔分布: 不再根据偏离程度为不同的增强数据分配不同的贝塔分布, 而是从一个固定的贝塔分布中采样混合权重. (5) 移除丰富模块: 移除 REDRec 的丰富模块. (6) 移除自适应聚合: 移除丰富模块中的自适应聚合过程, 转而使用平均聚合. (7) 移除两阶段训练: 移除第 4.3 节中所提出的两阶段训练策略, 转而在模型训练开始时就启用 REDRec. 消融实验的实验结果在表 7 中展示.

表 7 以 SASRec 为基础模型的消融实验结果

变体	Beauty		Sports		Yelp		Home	
	H@10	N@10	H@10	N@10	H@10	N@10	H@10	N@10
REDRec	0.0792	0.0452	0.0462	0.0257	0.0373	0.0192	0.0247	0.0134
加权融合损失	0.0784	0.0445	0.0450	0.0243	0.0362	0.0183	0.0245	0.0138
硬标签主导的互补	0.0779	0.0442	0.0456	0.0240	0.0379	0.0190	0.0234	0.0128
移除修复模块	0.0684	0.0388	0.0388	0.0206	0.0329	0.0166	0.0220	0.0114
固定贝塔分布	0.0732	0.0421	0.0440	0.0241	0.0341	0.0171	0.0231	0.0123
移除丰富模块	0.0725	0.0411	0.0415	0.0220	0.0348	0.0177	0.0225	0.0118
移除自适应聚合	0.0783	0.0438	0.0452	0.0248	0.0356	0.0180	0.0236	0.0126
移除两阶段训练	0.0780	0.0434	0.0450	0.0241	0.0364	0.0185	0.0233	0.0124

从表 7 中可以观察到, 当改用其他损失融合策略时, 模型性能并没有出现明显上升. 同时, REDRec 中的所有模块和设计均对最终性能有所贡献. 无论是采用加权融合损失或是硬标签主导的互补损失, 模型的最终性能没有出现明显的提高, 在多数案例中弱于不使用任何的加权的损失函数. 这表明同时均等地使用 3 个学习目标可以更充分地利用来自两个模块的所有增强样本, 以达到更好的推荐性能. 当移除修复模块时, 模型的推荐准确性下降最为显著, 这进一步验证了现有启发式数据增强方法存在的语义信息和偏好知识破坏问题, 同时也凸显了 REDRec 的修复效果. 同时, 丰富模块则为训练过程提供了更为丰富和与修复后嵌入表示相匹配的标签信息, 移除它之后模型性能也出现了下降. 当采用固定的贝塔分布为所有待修复的增强序列嵌入表示采样混合权重时, 一些偏离程度较高的序列可能无法得到妥善修复, 而一些偏离程度较低的序列中所蕴含的多样变化则可能被原始序列的嵌入表示所稀释. 此时, 修复效果欠佳, 而模型性能也出现了对应的下降. 当移除丰富模块的自适应聚合时, 一些包含噪声

或模型置信度较低的标签在最终标签中占比提高, 模型可能从这样的软标签中学习到的不准确的偏好知识. 此外, 若移除两阶段的训练策略, 在训练开始时就启用修复模块和丰富模块表示混杂, 则可能在嵌入表示混合和生成软标签时引入额外噪声或增加模型的收敛难度. 特别地, 对于丰富模块, 在模型训练初期骨干网络表征能力尚弱时, 预测得到的标签可能包含大量噪声和不准确的预测结果, 利用这些标签进行后续增强和训练则可能使模型学习到不准确甚至是错误的用户偏好.

5.7 场景泛化

为了进一步验证 REDRec 在不同推荐场景下的泛化性, 本文选取了一个短视频推荐数据集 (Douyin). 这一数据集中包含了用户的短视频观看记录. 它包含 20 398 位用户与 8 299 个短视频的 139 834 条交互记录. 平均序列长度为 6.9, 稀疏度为 99.92%. 这一平均序列长度低于已有的 4 个数据集, 将进一步验证 REDRec 在短序列上的性能表现. 从表 8 中可以看出, REDRec 在序列较短的短视频推荐场景下依然带来了显著的推荐性能提升, 进一步证明了其有效性与泛化性. 此外, 可以观察到, 一些启发式方法依然带来了性能损失 (例如 CMR 应用于 SASRec 模型), 与实证研究中的观察相似. 这也说明了偏好及语义信息受损问题的广泛性及对其进行修复的必要性.

表 8 Douyin 数据集上原始模型、添加启发式数据增强方法和进一步添加 REDRec 后的性能对比

方法	GRU4Rec		NextItNet		SASRec		FMLP-Rec		LRURec	
	$H@10$	$N@10$	$H@10$	$N@10$	$H@10$	$N@10$	$H@10$	$N@10$	$H@10$	$N@10$
基础模型	0.1060	0.0572	0.0696	0.0372	0.1243	0.0739	0.1245	0.0706	0.1239	0.0742
+CMR	0.1103	0.0635	0.0782	0.0403	0.1174	0.0705	0.1335	0.0742	0.1362	0.0771
w/ REDRec-CMR	0.1312	0.0779	0.0976	0.0594	0.1608	0.1054	0.1766	0.1098	0.1805	0.1167
+CMRSI	0.1129	0.0715	0.0825	0.0437	0.1423	0.0916	0.1441	0.0923	0.1409	0.0797
w/ REDRec-CMRSI	0.1429	0.0892	0.1124	0.0715	0.1735	0.1083	0.1936	0.1185	0.1873	0.1180

注: w/ REDRec-CMR和w/ REDRec-CMRSI表示在相应启发式数据增强方法上添加REDRec. 加粗代表REDRec方法在启发式数据增强方法的基础上带来了性能提升

5.8 超参数敏感性分析

敏感性分析实验调查了过滤百分比 Q 、用于初始化贝塔分布的 α_a 和 α_b 以及区间数量 T 对模型性能的影响. 实验按照第 5.4 节中的范围调整这 3 个超参数, 观察推荐性能的变化情况.

图 4 展示了过滤百分比 Q 的大小对模型性能的影响. 这一超参数决定了有多少偏离程度过高的序列在启发式数据增强之后被直接过滤掉, 而不会参与接下来的修复和丰富过程. 从图中可以看到, 当 Q 从最小值 5% 逐渐上升到 10% 时, 模型性能在 3 个数据集上都取得了提升, 这说明提高 Q 进一步过滤掉了更多在语义信息和偏好知识上过于偏离的序列, 修复和丰富剩余序列有助于提高模型性能. 但是, 当 Q 进一步上升时, 模型性能出现了下降. 在 Yelp 数据集上, 模型性能 Q 在大于 20% 时也出现了下降. 这一现象说明, 当 Q 的值过大时, 原本值得修复和丰富的增强序列也在这一过程中被过滤掉了. 这些序列虽然在一定程度上存在语义信息和偏好知识受损的情况, 但经过修复和丰富后, 依然可以作为高质量的样本参与训练. 因此, 保留它们是有价值的. 在实践中, 推荐使用者优先在 10%–20% 这一区间内调整 Q 的取值.

图 5 展示了用于初始化贝塔分布的 α_a 和 α_b 的值对模型性能的影响. 用橘色高亮了性能最优的设置. 例如, 12-2 表示初始 $\alpha_a = 12$ 且 $\alpha_b = 2$. 可以观察到, α_a 和 α_b 的最优取值往往偏向于整体较小但差异较大的组合. 整体偏小使得每个贝塔分布的概率密度函数更为平滑, 减少了极端 λ 值 (接近 0 或 1) 出现的情况. 这避免了修复后的样本中原始序列的嵌入表示或增强序列的嵌入表示其中一方占据主导的情况. 同时 α_a 和 α_b 较大插值使得 T 个贝塔分布的差异更为明显, 每个不同的贝塔分布都可以为对应偏离程度的增强序列提供独特的混合权重, 修复受损的语义信息和偏好知识. 此外, 过小或过大的 α_a 和 α_b 组合值都无法取得较好的性能. 前者使贝塔分布的概率密度函数过于平滑, 即 λ 偏向于 0.5, 削弱了修复模块为不同偏离程度样本提供不同混合权重的能力. 而后者则加剧了极端 λ 值 (接近 0 或 1) 出现的情况, 使得混合样本中可能有一方占据绝对主导. 在实践中, 推荐使用者优先在 14-2 和 14-4 中选择 α_a 和 α_b 的初始值.

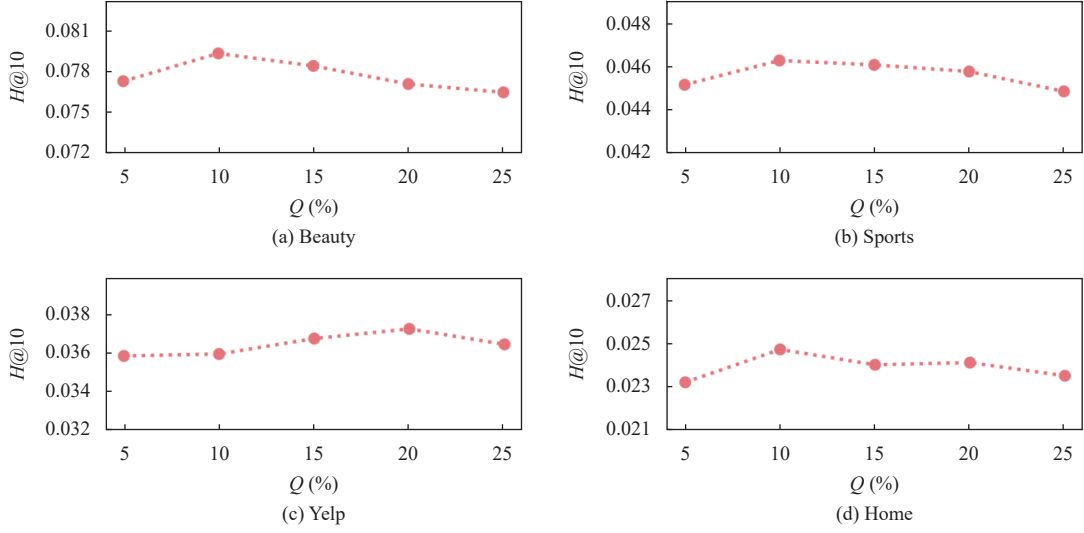
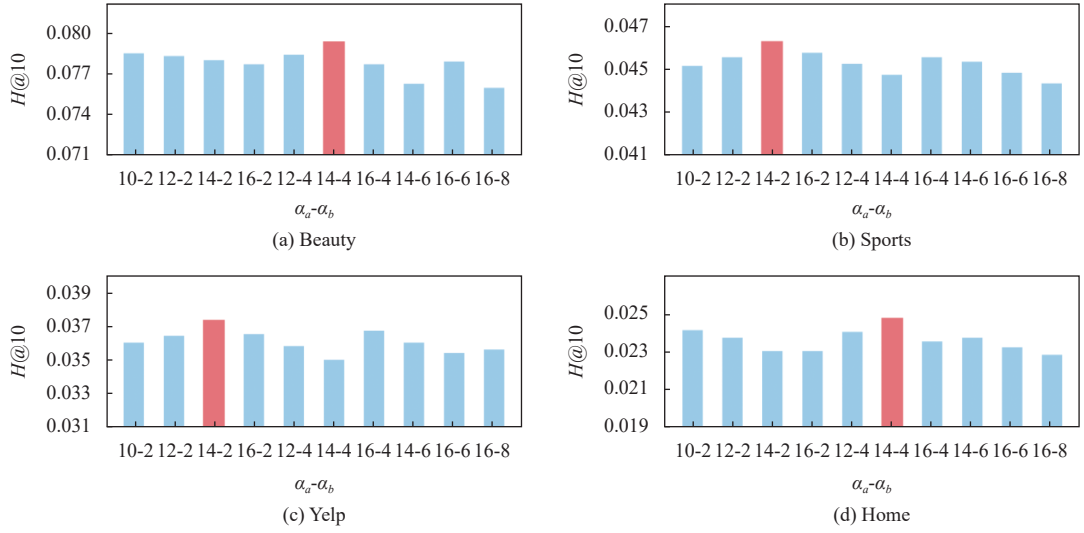
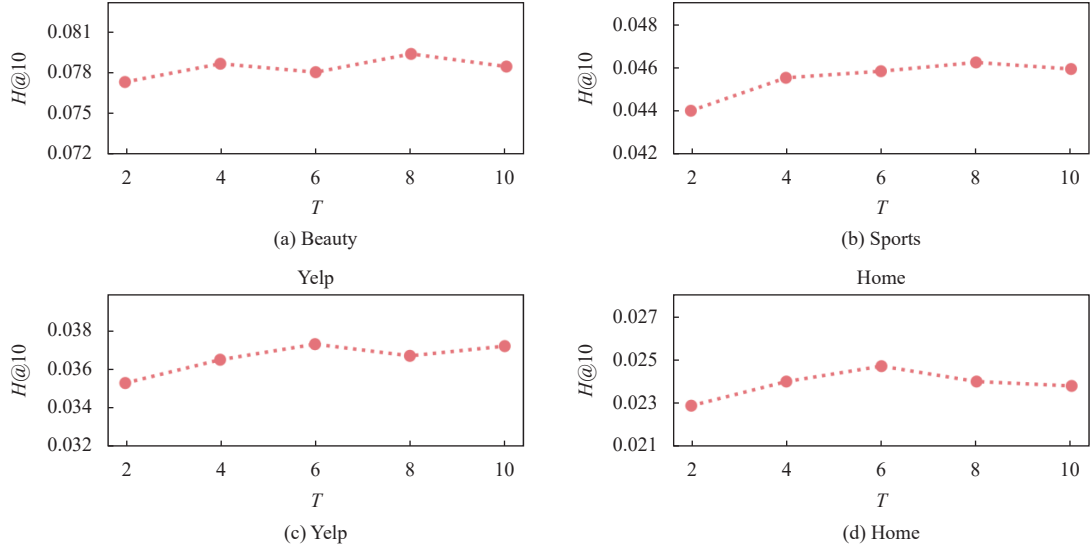
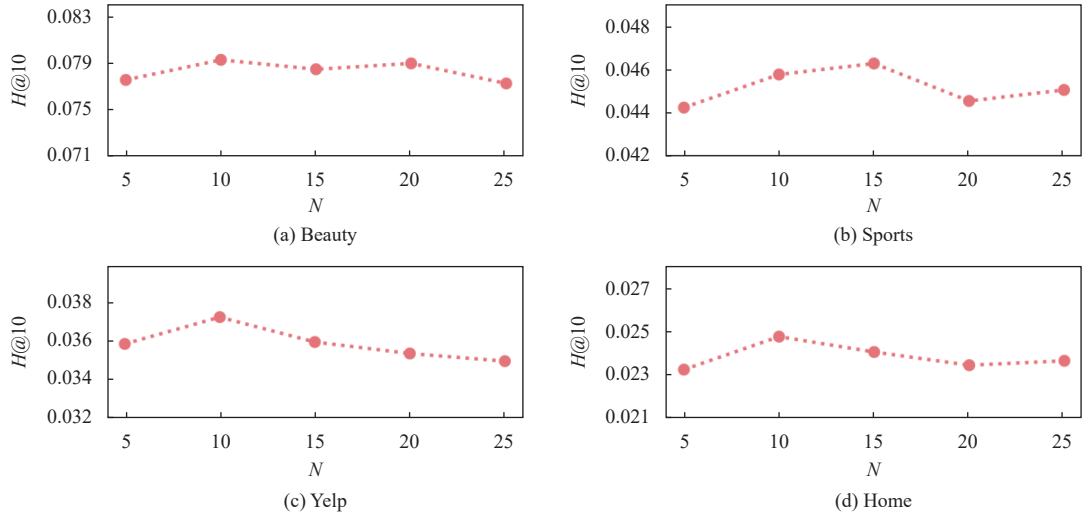
图4 过滤百分比 Q 对模型性能的影响图5 初始化贝塔分布的 α_a 和 α_b 对模型性能的影响.

图6展示了区间数量 T 对模型性能的影响. 这一参数决定了修复模块根据初始 α_a 和 α_b 的值生成多少个不同的贝塔分布. 可以观察到, 随着 T 的取值从2逐渐提高, 模型的性能也随之上升. 这是因为修复模块可以根据增强数据的偏离和受损程度进行更细粒度的修复, 不同偏离程度的序列都有对应的混合权重为之分配, 进而达到更令人满意的修复效果. 当 T 的取值进一步上升时, 模型的性能不会进一步提高. 此外, 可以观察到, T 的最优取值往往是6或8. 因此, 在实践中, 推荐使用者优先尝试这两个数值. 在未来, 也计划进一步改进修复模块, 使其可以根据增强数据的情况 (如平均序列长度和稀疏长度) 自适应地调整和搜索最优的 α_a 、 α_b 和 T 的取值.

图7展示了扰动次数 N 对模型性能的影响. 这一超参数决定了对输出表示 H_u^r 进行扰动得到扰动表示的数量. 从图中可以观察到, 在扰动次数由5上升到10时, 模型性能出现了明显上升. 这是因为随着扰动序列的数量增多, 聚合得到的表示可以更好地模拟用户行为的不确定性并促进模型捕捉用户的核心偏好. 然而, 当扰动次数进一步增加时, 模型性能均出现了下降, 这可能是由于扰动次数过多而引入了噪声, 模型难以从聚合得到的表示中提取到准确的用户行为模式, 进而导致推荐准确性受损.

图6 区间数量 T 对模型性能的影响图7 扰动次数 N 对模型性能的影响

5.9 REDRec 修复效果的分析

为了分析不同方法产生的增强序列与原始序列的语义偏离程度, 实验计算了不同增强方法产生的增强序列与原始序列嵌入的詹森-香农散度. 为了保持分析的公平性, 本文首先利用原始数据训练一个 SASRec 模型, 待损失充分收敛后, 利用训练好的嵌入表得到所需序列数据的嵌入表示, 随后计算所有序列的平均詹森-香农散度. 从表 9 中可以观察到, 启发式数据增强方法产生的增强序列与原始序列在语义和偏好信息上的偏离最大, 模型因此无法从这些数据中学习准确的偏好. 特别是 Reorder 的重排序操作可能破坏了原本的序列模式, 使增强序列在很大程度上偏离原始序列 (詹森-香农散度). 基于模型的增强方法, 如 DR4SR 和 DiffuASR, 以及精心设计的嵌入层面增强方法 BSARec, 这些方法在一定程度上保留了原始序列的语义和偏好信息, 因此它们产生的增强数据和原始数据的关联性较强. 对于本文所提出的 REDRec, 其修复模块有效保留了原始序列中重要语义信息 (詹森-香农散度), 使增强序列在包含多样性变化的同时又包含了原始序列的偏好知识, 模型能够从中学习到准确的用户偏

好并给出令人满意的推荐结果.

表 9 不同增强方法产生的增强序列与原始序列嵌入的平均詹森-香农散度

方法	Beauty	Sports	Yelp	Home	平均值
Crop	0.2768	0.3017	0.2136	0.2235	0.2539
Mask	0.3784	0.2816	0.1713	0.2451	0.2691
Reorder	0.1453	0.1965	0.1365	0.1978	0.1690
Substitute	0.4292	0.3555	0.2367	0.2712	0.3232
Insert	0.3925	0.3787	0.2545	0.2903	0.3290
CMR	0.3042	0.2760	0.1843	0.2286	0.2483
CMRSI	0.3365	0.3469	0.2109	0.2508	0.2863
DR4SR	0.4541	0.4302	0.2580	0.3225	0.3662
DiffuASR	0.3807	0.4642	0.2654	0.2913	0.3504
BSARec	0.4735	0.5028	0.3081	0.3623	0.4117
REDRec	0.5157	0.5615	0.3479	0.4114	0.4591

此外, 采用第 3 节实证研究的实验设置和可视化方式, 对原始序列及经启发式算子增强并由 REDRec 修复后的序列进行可视化呈现. 结果如图 8 所示. 此处所选择的用户序列与图 2 中所选择的用户序列相同. 通过对比图 2 与图 8 可以发现, 经过本文所提出的 REDRec 修复后的序列的嵌入表示在整体分布上更接近原始序列的嵌入表示, 有效保留了重要的语义信息和用户偏好知识. 与此同时, 原始序列表示与修复后表示之间的适度差异则进一步拓展了训练样本空间, 使模型能够学习到更全面的用户偏好, 从而进一步提高自身的推荐准确性和泛化性.

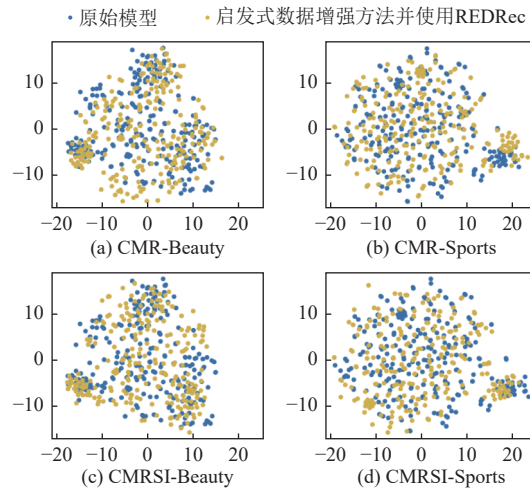


图 8 原始序列嵌入表示及经过 REDRec 修复后的增强序列嵌入表示的可视化呈现

6 总 结

本研究揭示了现有启发式数据增强方法对原始序列中语义信息和偏好知识的破坏, 以及标签的稀缺与不匹配问题. 实证研究表明, 现有启发式数据增强方法生成的增强数据分布与原始数据存在显著偏差, 导致模型学习到不准确甚至是错误的用户偏好. 同时, 简单的詹森-香农散度过滤策略即可缓解这一问题. 基于已有问题和实证研究中的发现, 本文提出了一种用于修复和丰富序列推荐中启发式数据增强的框架 (REDRec). 通过精心设计的修复模块与丰富模块, REDRec 有效修复了被启发式数据增强算子所损害的语义信息和偏好知识, 同时生成了与修复后样本所匹配的高质量标签. 两个模块共同为模型提供了既保留原始序列的重要信息, 又包含丰富变化的高质量训练数据. 此外, 本文比较了 REDRec 与现有方法的优势. 通过全面的实验验证和分析, REDRec 展现出卓越的有效

性和强大的泛化能力. 此外, 本文还通过量化分析和可视化分析的方法展示了 REDRec 相较于其他方法对受损增强数据的修复能力. REDRec 使增强序列在包含多样性变化的同时又包含了原始序列的偏好知识, 模型能够从中学到准确的用户偏好并给出令人满意的推荐结果. 在未来的研究中, 将借助物品描述文本和图片等模态信息, 以探索更优的语义差异衡量方法. 此外, 还将继续扩展和改进 REDRec, 使其能够被用于基于模型的序列增强方法或其他推荐领域的的数据增强方法.

References

- [1] Pan LW, Pan WK, Wei MY, Yin HZ, Ming Z. A survey on sequential recommendation. *Frontiers of Computer Science*, 2026, 20(3): 2003606. [doi: [10.1007/s11704-025-41329-w](https://doi.org/10.1007/s11704-025-41329-w)]
- [2] Sun WQ, Xie RB, Zhang JJ, Zhao X, Kang ZH, Wen JR. A method of lifelong sequential recommendation based on matrix and mixing dual-flow long short-term memory networks. *Chinese Journal of Computers*, 2025, 48(12): 2789–2808 (in Chinese with English abstract). [doi: [10.11897/SP.J.1016.2025.02789](https://doi.org/10.11897/SP.J.1016.2025.02789)]
- [3] Kang WC, McAuley J. Self-attentive sequential recommendation. In: *Proc. of the 2018 IEEE Int'l Conf. on Data Mining (ICDM)*. Singapore: IEEE, 2018. 197–206. [doi: [10.1109/ICDM.2018.00035](https://doi.org/10.1109/ICDM.2018.00035)]
- [4] Zhou K, Yu H, Zhao WX, Wen JR. Filter-enhanced MLP is all you need for sequential recommendation. In: *Proc. of the 2022 ACM Web Conf.* Lyon: ACM, 2022. 2388–2399. [doi: [10.1145/3485447.3512111](https://doi.org/10.1145/3485447.3512111)]
- [5] Lu XK, Feng J, Han YQ, Wang H, Chen EH. GraphMLP-mixer: A graph-MLP architecture for efficient multi-behavior sequential recommendation method. *Journal of Computer Research and Development*, 2024, 61(8): 1917–1929 (in Chinese with English abstract). [doi: [10.7544/issn1000-1239.202440137](https://doi.org/10.7544/issn1000-1239.202440137)]
- [6] Zang TZ, Zhu YM, Liu HB, Zhang RH, Yu JD. A survey on cross-domain recommendation: Taxonomies, methods, and future directions. *ACM Trans. on Information Systems*, 2023, 41(2): 42. [doi: [10.1145/3548455](https://doi.org/10.1145/3548455)]
- [7] Qian ZS, Wan ZL, Fan FY, Fu TF. Dual-view self-supervised session-based recommendation model based on temporal interval aware data augmentation. *Ruan Jian Xue Bao/Journal of Software*, 2025, 36(12): 5695–5719 (in Chinese with English abstract). <http://www.jos.org.cn/1000-9825/7404.htm> [doi: [10.13328/j.cnki.jos.007404](https://doi.org/10.13328/j.cnki.jos.007404)]
- [8] Dang YZ, Liu YT, Yang EN, Guo GB, Jiang LY, Zhao JZ, Wang XW. Efficient and adaptive recommendation unlearning: A guided filtering framework to erase outdated preferences. *ACM Trans. on Information Systems*, 2025, 43(2): 49. [doi: [10.1145/3706633](https://doi.org/10.1145/3706633)]
- [9] Wang ZT, Wang PF, Liu KP, Wang PY, Fu YJ, Lu CT, Aggarwal CC, Pei J, Zhou YC. A comprehensive survey on data augmentation. *IEEE Trans. on Knowledge and Data Engineering*, 2026, 38(1): 47–66. [doi: [10.1109/TKDE.2025.3622600](https://doi.org/10.1109/TKDE.2025.3622600)]
- [10] Yang SR, Xiao WK, Zhang MC, Guo SH, Zhao J, Shen FR. Image data augmentation for deep learning: A survey. *arXiv:2204.08610*, 2023.
- [11] Joshi S, Feng YS, Hsiao KJ, Zhang Z, Lamkhede S. Sliding window training-utilizing historical recommender systems data for foundation models. In: *Proc. of the 18th ACM Conf. on Recommender Systems*. Bari: ACM, 2024. 835–837. [doi: [10.1145/3640457.3688051](https://doi.org/10.1145/3640457.3688051)]
- [12] Xie X, Sun F, Liu ZY, Wu SW, Gao JY, Zhang JD, Ding BL, Cui B. Contrastive learning for sequential recommendation. In: *Proc. of the 38th IEEE Int'l Conf. on Data Engineering (ICDE)*. Kuala Lumpur: IEEE, 2022. 1259–1273. [doi: [10.1109/ICDE53745.2022.00099](https://doi.org/10.1109/ICDE53745.2022.00099)]
- [13] Liu ZW, Chen YJ, Li J, Yu PS, McAuley J, Xiong CM. Contrastive self-supervised sequential recommendation with robust augmentation. *arXiv:2108.06479*, 2021.
- [14] Jiang JY, Zhang PY, Luo YT, Li CZ, Kim JB, Zhang K, Wang SZ, Kim S, Yu PS. Improving sequential recommendations via bidirectional temporal data augmentation with pre-training. *IEEE Trans. on Knowledge and Data Engineering*, 2025, 37(5): 2652–2664. [doi: [10.1109/TKDE.2025.3546035](https://doi.org/10.1109/TKDE.2025.3546035)]
- [15] Liu QD, Yan F, Zhao XY, Du ZC, Guo HF, Tang RM, Tian F. Diffusion augmentation for sequential recommendation. In: *Proc. of the 32nd ACM Int'l Conf. on Information and Knowledge Management*. Birmingham: ACM, 2023. 1576–1586. [doi: [10.1145/3583780.3615134](https://doi.org/10.1145/3583780.3615134)]
- [16] Dang YZ, Yang EN, Guo GB, Jiang LY, Wang XW, Xu XX, Sun QH, Liu H. TiCoSeRec: Augmenting data to uniform sequences by time intervals for effective recommendation. *IEEE Trans. on Knowledge and Data Engineering*, 2024, 36(6): 2686–2700. [doi: [10.1109/tkde.2023.3324312](https://doi.org/10.1109/tkde.2023.3324312)]
- [17] Dang YZ, Yang EN, Liu YT, Guo GB, Jiang LY, Zhao JZ, Wang XW. Data augmentation for sequential recommendation: A survey. *arXiv:2409.13545*, 2024.
- [18] Wang DD, Liu HY, Liu Z, Yuan JD. Self-supervised recommendation method based on multiple views. *Ruan Jian Xue Bao/Journal of*

- Software, 2026, 37(2): 684–699 (in Chinese with English abstract). <http://www.jos.org.cn/1000-9825/7419.htm> [doi: 10.13328/j.cnki.jos.007419]
- [19] Karras T, Aittala M, Hellsten J, Laine S, Lehtinen J, Aila T. Training generative adversarial networks with limited data. In: Proc. of the 34th Int'l Conf. on Neural Information Processing Systems. Vancouver: Curran Associates Inc., 2020. 1015.
 - [20] Zhang HY, Cisse M, Dauphin YN, Lopez-Paz D. mixup: Beyond empirical risk minimization. In: Proc. of the 6th Int'l Conf. on Learning Representations. Vancouver: OpenReview.net, 2018.
 - [21] He RN, McAuley J. Fusing similarity models with Markov chains for sparse sequential recommendation. In: Proc. of the 16th IEEE Int'l Conf. on Data Mining (ICDM). Barcelona: IEEE, 2016. 191–200. [doi: 10.1109/ICDM.2016.0030]
 - [22] Hidasi B, Karatzoglou A, Baltrunas L, Tikk D. Session-based recommendations with recurrent neural networks. In: Proc. of the 4th Int'l Conf. on Learning Representations. San Juan: OpenReview.net, 2016. [doi: 10.48550/arXiv.1511.06939]
 - [23] Cui Q, Wu S, Liu Q, Zhong W, Wang L. MV-RNN: A multi-view recurrent neural network for sequential recommendation. IEEE Trans. on Knowledge and Data Engineering, 2020, 32(2): 317–331. [doi: 10.1109/TKDE.2018.2881260]
 - [24] Qu SL, Yuan FJ, Guo GB, Zhang LG, Wei W. CmnRec: Sequential recommendations with chunk-accelerated memory network. IEEE Trans. on Knowledge and Data Engineering, 2023, 35(4): 3540–3550. [doi: 10.1109/TKDE.2022.3141102]
 - [25] Liu XY, Yin HL, Zang YL, Wu WL, Zhuo JN, Xu FR, Chen LY, Ma WH, Li BH. A graph-based interpolation sequential recommender with deformable convolutional network. Journal of Computer Research and Development, 2025, 62(10): 2583–2594 (in Chinese with English abstract). [doi: 10.7544/issn1000-1239.202220749]
 - [26] Tang JX, Wang K. Personalized Top-N sequential recommendation via convolutional sequence embedding. In: Proc. of the 11th ACM Int'l Conf. on Web Search and Data Mining. Marina: ACM, 2018. 565–573. [doi: 10.1145/3159652.3159656]
 - [27] Yuan FJ, Karatzoglou A, Arapakis I, Jose JM, He XN. A simple convolutional generative network for next item recommendation. In: Proc. of the 12th ACM Int'l Conf. on Web Search and Data Mining. Melbourne: ACM, 2019. 582–590. [doi: 10.1145/3289600.3290975]
 - [28] Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, Kaiser Ł, Polosukhin I. Attention is all you need. In: Proc. of the 31st Int'l Conf. on Neural Information Processing Systems. Long Beach: Curran Associates Inc., 2017. 6000–6010.
 - [29] Li JC, Wang YJ, McAuley J. Time interval aware self-attention for sequential recommendation. In: Proc. of the 13th Int'l Conf. on Web Search and Data Mining. Houston: ACM, 2020. 322–330. [doi: 10.1145/3336191.3371786]
 - [30] Fan ZW, Liu ZW, Wang Y, Wang A, Nazari Z, Zheng L, Peng H, Yu PS. Sequential recommendation via stochastic self-attention. In: Proc. of the 2022 ACM Web Conf. Lyon: ACM, 2022. 2036–2047. [doi: 10.1145/3485447.3512077]
 - [31] Zhao RM, Sun SY, Yan FL, Peng J, Ju SG. Multi-interest aware sequential recommender system based on contrastive learning. Journal of Computer Research and Development, 2024, 61(7): 1730–1740 (in Chinese with English abstract). [doi: 10.7544/issn1000-1239.202330622]
 - [32] Yue ZR, Wang YQ, He ZK, Zeng HM, McAuley J, Wang D. Linear recurrent units for sequential recommendation. In: Proc. of the 17th ACM Int'l Conf. on Web Search and Data Mining. Merida: ACM, 2024. 930–938. [doi: 10.1145/3616855.3635760]
 - [33] Tan YK, Xu XX, Liu Y. Improved recurrent neural networks for session-based recommendations. In: Proc. of the 1st Workshop on Deep Learning for Recommender Systems. Boston: ACM, 2016. 17–22. [doi: 10.1145/2988450.2988452]
 - [34] Chen YJ, Liu ZW, Li J, McAuley J, Xiong CM. Intent contrastive learning for sequential recommendation. In: Proc. of the 2022 ACM Web Conf. Lyon: ACM, 2022. 2172–2182. [doi: 10.1145/3485447.3512090]
 - [35] Dang YZ, Zhang JH, Liu YT, Yang EN, Liang YL, Guo GB, Zhao JZ, Wang XW. Augmenting sequential recommendation with balanced relevance and diversity. In: Proc. of the 39th AAAI Conf. on Artificial Intelligence. Philadelphia: AAAI, 2025. 11563–11571. [doi: 10.1609/aaai.v39i11.33258]
 - [36] Hu SR, Wu WC, Tang ZL, Huan ZX, Wang L, Zhang XL, Zhou J, Zou LX, Li CL. HORAE: Temporal multi-interest pre-training for sequential recommendation. ACM Trans. on Information Systems, 2025, 43(4): 88. [doi: 10.1145/3727645]
 - [37] Xiao S, Zhang JT, Tang CG, Huang ZZ. Frequency-domain disentanglement-fusion and dual contrastive learning for sequential recommendation. In: Proc. of the 34th ACM Int'l Conf. on Information and Knowledge Management. Seoul: ACM, 2025. 3498–3508. [doi: 10.1145/3746252.3761220]
 - [38] Liu ZL, Xu Y, Cong G, Zhu L, Qiu QJ, Zhang HX. ARTS: A general and efficient multi-task self-prompt framework for explainable sequential recommendation. ACM Trans. on Information Systems, 2025, 43(3): 73. [doi: 10.1145/3717833]
 - [39] Wang ZL, Zhang JS, Xu HT, Chen X, Zhang YF, Zhao WX, Wen JR. Counterfactual data-augmented sequential recommendation. In: Proc. of the 44th Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval. ACM, 2021. 347–356. [doi: 10.1145/3404835.3462855]
 - [40] Liu ZW, Fan ZW, Wang Y, Yu PS. Augmenting sequential recommendation with pseudo-prior items via reversely pre-training

- Transformer. In: Proc. of the 44th Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval. ACM, 2021. 1608–1612. [doi: [10.1145/3404835.3463036](https://doi.org/10.1145/3404835.3463036)]
- [41] Wang JL, Le Y, Chang B, Wang YY, Chi EH, Chen MM. Learning to augment for casual user recommendation. In: Proc. of the 2022 ACM Web Conf. Lyon: ACM, 2022. 2183–2194. [doi: [10.1145/3485447.3512147](https://doi.org/10.1145/3485447.3512147)]
- [42] Yin MJ, Wang H, Guo W, Liu Y, Zhang SJ, Zhao SR, Lian DF, Chen EH. Dataset regeneration for sequential recommendation. In: Proc. of the 30th ACM SIGKDD Conf. on Knowledge Discovery and Data Mining. Barcelona: ACM, 2024. 3954–3965. [doi: [10.1145/3637528.3671841](https://doi.org/10.1145/3637528.3671841)]
- [43] Zhou R, Wu X, Qiu ZP, Zheng YF, Chen X. Distributionally robust sequential recommendation. In: Proc. of the 46th Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval. Taipei: ACM, 2023. 279–288. [doi: [10.1145/3539618.3591668](https://doi.org/10.1145/3539618.3591668)]
- [44] Liao YX, Yang YH, Hou M, Wu L, Xu HF, Liu H. Mitigating distribution shifts in sequential recommendation: An invariance perspective. In: Proc. of the 48th Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval. Padua: ACM, 2025. 1603–1613. [doi: [10.1145/3726302.3730036](https://doi.org/10.1145/3726302.3730036)]
- [45] Lin XL, Pan WK, Ming Z. Towards interest drift-driven user representation learning in sequential recommendation. In: Proc. of the 48th Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval. Padua: ACM, 2025. 1541–1551. [doi: [10.1145/3726302.3730099](https://doi.org/10.1145/3726302.3730099)]
- [46] Zhang KK, Cao Q, Sun F, Liu XR, Shen HW, Cheng XQ. AsarRec: Adaptive sequential augmentation for robust self-supervised sequential recommendation. arXiv:2512.14047, 2026.
- [47] McAuley J, Pandey R, Leskovec J. Inferring networks of substitutable and complementary products. In: Proc. of the 21st ACM SIGKDD Int'l Conf. on Knowledge Discovery and Data Mining. Sydney: ACM, 2015. 785–794. [doi: [10.1145/2783258.2783381](https://doi.org/10.1145/2783258.2783381)]
- [48] van der Maaten L, Hinton G. Visualizing data using t-SNE. Journal of Machine Learning Research, 2008, 9(86): 2579–2605.
- [49] Wang Y, Zhang HR, Liu ZW, Yang LW, Yu PS. ContrastVAE: Contrastive variational autoencoder for sequential recommendation. In: Proc. of the 31st ACM Int'l Conf. on Information & Knowledge Management. Atlanta: ACM, 2022. 2056–2066. [doi: [10.1145/3511808.3557268](https://doi.org/10.1145/3511808.3557268)]
- [50] Krichene W, Rendle S. On sampled metrics for item recommendation. In: Proc. of the 26th ACM SIGKDD Int'l Conf. on Knowledge Discovery & Data Mining. ACM, 2020. 1748–1757. [doi: [10.1145/3394486.3403226](https://doi.org/10.1145/3394486.3403226)]
- [51] Kingma DP, Ba J. Adam: A method for stochastic optimization. In: Proc. of the 3rd Int'l Conf. on Learning Representations. San Diego: OpenReview.net, 2015. [doi: [10.48550/arXiv.1412.6980](https://doi.org/10.48550/arXiv.1412.6980)]

附中文参考文献

- [2] 孙文奇, 谢若冰, 张君杰, 赵鑫, 康战辉, 文继荣. 基于矩阵与混合机制双流长短期记忆网络的终身长序列推荐方法. 计算机学报, 2025, 48(12): 2789–2808. [doi: [10.11897/SP.J.1016.2025.02789](https://doi.org/10.11897/SP.J.1016.2025.02789)]
- [5] 卢晓凯, 封军, 韩永强, 王皓, 陈恩红. GraphMLP-Mixer: 基于图-多层感知机架构的高效多行为序列推荐方法. 计算机研究与发展, 2024, 61(8): 1917–1929. [doi: [10.7544/issn1000-1239.202440137](https://doi.org/10.7544/issn1000-1239.202440137)]
- [7] 钱忠胜, 万子珑, 范赋宇, 付庭峰. 结合时间间隔数据增强的对偶视图自监督会话推荐模型. 软件学报, 2025, 36(12): 5695–5719. <http://www.jos.org.cn/1000-9825/7404.htm> [doi: [10.13328/j.cnki.jos.007404](https://doi.org/10.13328/j.cnki.jos.007404)]
- [18] 王丹丹, 刘海洋, 刘壮, 原继东. 基于多视角的自监督推荐方法. 软件学报, 2026, 37(2): 684–699. <http://www.jos.org.cn/1000-9825/7419.htm> [doi: [10.13328/j.cnki.jos.007419](https://doi.org/10.13328/j.cnki.jos.007419)]
- [25] 刘昕悦, 尹海莲, 臧亚磊, 吴文隆, 卓俊男, 徐凤如, 陈吕莹, 马维华, 李博涵. 基于图插值和可变形卷积网络的序列推荐. 计算机研究与发展, 2025, 62(10): 2583–2594. [doi: [10.7544/issn1000-1239.202220749](https://doi.org/10.7544/issn1000-1239.202220749)]
- [31] 赵容梅, 孙思雨, 鄢凡力, 彭舰, 琚生根. 基于对比学习的多兴趣感知序列推荐系统. 计算机研究与发展, 2024, 61(7): 1730–1740. [doi: [10.7544/issn1000-1239.202330622](https://doi.org/10.7544/issn1000-1239.202330622)]

作者简介

党翌洲, 博士生, CCF 学生会员, 主要研究领域为推荐系统, 数据挖掘.

赵楚, 博士生, 主要研究领域为推荐系统, 大语言模型, 数据挖掘.

姜琳颖, 硕士, 副教授, 主要研究领域为嵌入式系统, 数据分析.

马连博, 博士, 教授, 博士生导师, CCF 杰出会员, 主要研究领域为工业互联网, 多目标优化, 联邦学习.

郭贵冰, 博士, 教授, 博士生导师, CCF 专业会员, 主要研究领域为推荐系统, 自然语言处理, 数据挖掘.