

结合时间间隔数据增强的对偶视图自监督会话推荐模型^{*}

钱忠胜, 万子珑, 范赋宇, 付庭峰

(江西财经大学 计算机与人工智能学院, 江西 南昌 330013)

通信作者: 钱忠胜, E-mail: changesme@163.com



摘 要: 会话推荐旨在基于用户的一系列项目预测其交互的下一项目, 现有大多数会话推荐对于会话内项目间的时间间隔信息利用不够充分, 影响推荐准确性. 近年, 图神经网络凭借自身强大的复杂关系建模能力在会话推荐中受到推崇, 但仅基于图神经网络的会话推荐忽略了会话间的隐藏高阶关系, 信息不够丰富. 此外, 数据稀疏性一直是推荐系统中存在的现象, 研究中多使用对比学习对此实施改善, 然而大多对比学习框架形式单一, 泛化能力不强. 基于此, 提出一种结合自监督学习的会话推荐模型. 首先, 该模型利用用户会话内项目间的时间间隔信息对会话序列实施数据增强, 丰富会话内信息, 以提高推荐准确性; 其次, 构建超图卷积网络和 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, 以丰富推荐多样性; 最后, 融合数据增强后的会话内信息、多视角下的会话间信息以及原始会话信息进行对比学习, 以增强模型泛化性. 通过与 11 个已有经典模型在 4 个数据集上的对比发现, 所提模型是可行高效的, 在 HR 与 $NDCG$ 指标上分别平均提升 5.96%、5.89%.

关键词: 会话推荐; 自监督学习; 超图卷积网络; 对偶视图; 数据增强

中图法分类号: TP181

中文引用格式: 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>

英文引用格式: 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). <http://www.jos.org.cn/1000-9825/7404.htm>

Dual-view Self-supervised Session-based Recommendation Model Based on Temporal Interval Aware Data Augmentation

QIAN Zhong-Sheng, WAN Zi-Long, FAN Fu-Yu, FU Ting-Feng

(School of Computer and Artificial Intelligence, Jiangxi University of Finance and Economics, Nanchang 330013, China)

Abstract: Session-based recommendation aims to predict the next item a user will interact with based on a series of items. Most existing session-based recommender systems do not fully utilize the temporal interval information between items within a session, affecting the accuracy of recommendations. In recent years, graph neural networks have gained significant attention in session-based recommendation due to their strong ability to model complex relationships. However, session-based recommendations that rely solely on graph neural networks overlook the hidden high-order relationships between sessions, resulting in less rich information. In addition, data sparsity has always been a phenomenon in recommender systems, and contrastive learning is often employed to address this issue. However, most contrastive learning frameworks lack strong generalization capabilities due to their singular form. Based on this, a session-based recommendation model combined with self-supervised learning is proposed. First, the model utilizes the temporal interval information between items within user sessions to perform data augmentation, enriching the information within the sessions to improve recommendation accuracy. Second, a dual-view encoder is constructed, combining a hypergraph convolutional network encoder and a

^{*} 基金项目: 国家自然科学基金 (62262025); 江西省自然科学基金重点项目 (20224ACB202012); 赣鄱俊才支持计划-主要学科学术和技术带头人培养项目-领军人才 (学术类) (20243BCE51024)

收稿时间: 2024-08-02; 修改时间: 2024-11-08, 2024-12-14; 采用时间: 2025-01-29; jos 在线出版时间: 2025-06-18

CNKI 网络首发时间: 2025-06-19

Transformer encoder to capture the hidden high-order relationships between sessions from multiple perspectives, thus enhancing the diversity of recommendations. Finally, the model integrates the augmented intra-session information, the multi-viewed inter-session information, and the original session information for contrastive learning to strengthen the model's generalization ability. Comparisons with 11 existing classic models on 4 datasets show that the proposed model is feasible and efficient, with average improvements of 5.96% and 5.89% on *HR* and *NDCG* metrics, respectively.

Key words: session-based recommendation; self-supervised learning; hypergraph convolutional network; dual-view; data augmentation

1 引言

随着电子商务的快速发展,消费者的购物行为逐渐呈现出高频短时、目标明确的特征.在此背景下,会话推荐系统 (session-based recommender system, SBRS) 通过精准捕捉用户短期行为,提供即时个性化推荐,成为电商平台提升用户体验和转化率的关键. SBRS 的核心在于充分挖掘用户行为数据,从会话内和会话间两个维度分析,全面捕捉用户的兴趣模式. 会话内分析聚焦于单个会话中用户的行为序列,挖掘目标导向和时序特征,从而预测短期兴趣并提供即时推荐;而会话间分析则侧重于跨会话的潜在高阶关系,揭示用户或商品间的隐式关联,发现更广泛的需求模式. 两者的协同建模不仅强化了系统的实时性与泛化性,还能平衡短期行为预测与全局兴趣刻画,显著提升推荐效果,为用户和平台创造更多价值.

早期的 SBRS 研究^[1]主要关注于在会话内分析用户与项目的交互数据预测用户短期兴趣. 尽管这些早期技术在当时取得初步成功,但往往对复杂用户行为建模能力不足. 随着机器学习尤其是深度学习技术的飞速发展,诸如基于循环神经网络 (recurrent neural network, RNN)^[2]或图神经网络 (graph neural network, GNN)^[3]的序列化推荐模型,被广泛应用于捕捉会话中的兴趣偏好. 此外,在捕获序列中的复杂依赖关系方面,基于 Transformer 等更多先进技术的模型已被证明表现较出色. 然而,这些模型所利用的时间信息过于单调,未合理处理用户交互项目间长短不一的时间间隔信息,从而影响模型预测的准确性.

除会话内的信息外,会话间的高阶关系近来也被研究人员发现会显著影响推荐结果. 而图卷积网络 (graph convolutional network, GCN) 因其强大的关系提取能力被应用于建模会话间的高阶关系. 其中,超图卷积网络 (hypergraph convolution network, HGCN)^[4-7]常被研究人员用于建模会话间关系,但其仅从超图节点的角度建模,考虑的角度比较单一. Xia 等人^[8]根据超图节点与超边的联系对项目与会话的关系进行建模,并融入自监督学习最大化节点与超边的互信息. 但对于图上的高阶邻居信息的提取程度不够,使得会话嵌入的多样性不足. 特别地,用户在会话中通常会表现多样化的兴趣. 单一地依赖行为序列往往仅能反映显式偏好,而结合隐藏的高阶关系 (如相似用户的行为模式) 可揭示更深层次的偏好. 在实际应用场景中,用户的行为模式常受到多种因素的影响 (如上下文、时间、场景等),隐藏的高阶关系能综合会话内外的信息进行全局建模,为复杂场景下的个性化推荐提供有力支撑.

例 1: 用户 A 购买了吸尘器、电视机、洗碗机、燃气灶,用户 B 购买了投影仪、扫地机器人、烤箱、燃气灶. 对于这两个会话,除了最后一项购买的商品相同外,没有相同的商品,一般的会话推荐系统并不会把这两个会话联系起来. 但是,这两个会话中有相似的商品且最后购买的商品是一样的,我们认为这是会话中的隐藏高阶关系.

近年来,研究人员将数据增强 (data augmentation) 与对比学习 (contrastive learning)^[9-11]引入推荐领域,缓解了数据稀疏性问题,同时也有效地融合会话内与会话间信息. 然而,现有方法多依赖简单的随机采样或启发式策略生成正负样本,可能会导致关键会话信息的丢失,尤其在会话推荐系统中更为显著. 对于会话内和会话间的特征建模,传统对比学习方法的局限性尤为突出. 就会话内而言,简单的数据增强策略无法有效捕捉行为序列中的时序依赖和短期兴趣变化,从而影响推荐的即时性和准确性. 而就会话间而言,需更精细建模用户间的复杂高阶关系. 现有的对比学习结构通常单一,仅能捕获局部信息,难以揭示会话间潜在的全局上下文关联,这限制了模型对跨会话关系的提取. 此外,单一尺度的对比学习框架难以适应多样化的用户行为模式和复杂场景中的动态需求. 因此,进一步结合数据增强与对比学习,通过整合会话内的时序信息和会话间的高阶关联,可全面地提取用户兴趣特征,提升模型的泛化能力和推荐性能,为复杂推荐任务提供更具鲁棒性和适应性的解决方案.

通过以上分析,我们发现基于数据增强的会话推荐系统存在以下问题.

- 1) 对于会话内时间信息的挖掘不够充分, 未合理处理用户交互项目间长短不一的时间间隔信息.
- 2) 数据增强方法仅考虑相似项目, 对高阶邻居信息的提取程度不够, 忽略了会话间的隐藏高阶关系.
- 3) 大多融合会话内与会话间信息的对比学习方式过于单一, 易受到无关项目的影响, 模型的泛化能力受限.

针对上述 3 个方面的问题, 本文提出结合时间间隔数据增强的对偶视图自监督会话推荐模型 (dual-view self-supervised session recommendation model based on temporal interval aware data augmentation, TIDA-DSSR). 首先, 挖掘会话内时间信息, 利用时间间隔信息对原始会话实施数据增强; 然后, 针对会话间的隐藏高阶关系, 构建对偶视图编码器从多角度对增强后的会话与原始会话进行编码; 最后, 利用对比学习采集不同视图内的信息, 通过预测任务生成推荐列表. 主要工作与贡献如下.

- 1) 利用会话内时间间隔信息进行数据增强, 有针对性地生成对比学习信号, 预防丢失关键会话项目特征, 丰富会话内信息.
- 2) 构建含有超图神经网络的对偶视图编码器, 将数据增强后的会话序列构成超图, 从多个视角建模用户会话, 结合时间间隔数据增强模块以更好地捕捉会话间的隐藏高阶关系.
- 3) 构造自监督任务框架并将对比学习融入模型, 让会话内与会话间不同视图的编码器互相获取信息以提升模型的泛化能力, 缓解会话序列中数据稀疏性对推荐任务造成的影响.
- 4) 在 4 个公开数据集上展开了对比与消融实验, 与当前经典的、主流的基于数据增强的推荐模型相比, 本文模型 TIDA-DSSR 在性能上优势较明显, 也验证了模型中各组件存在的必要性以及组合的合理性.

2 相关工作

SBRS 依靠用户交互数据预测未来行为, 但常受限于数据稀疏性和噪声. 数据增强通过生成训练样本提升模型鲁棒性^[12]. 图神经网络挖掘用户物品间复杂关系, 提高推荐准确性. 本节综述数据增强与图神经网络两者在推荐中的应用及挑战.

● 基于数据增强的推荐方法. 该类模型是在原始用户会话的基础上进行一系列变换得到信息更丰富的会话序列, 并基于此进行建模的方法, 但其获得的增强数据质量不如原始数据. 为解决该问题, 许多研究人员开始使用自监督学习^[13]让增强数据与原始数据进行相似度对比. Xie 等人^[9]提出 3 种用于预训练推荐任务的数据增强算子 (Crop, Reorder, Mask), 并配备了相应的对比学习框架. Yao 等人^[14]提出一种两阶段增强策略, 先掩盖项目嵌入层, 再放弃对比学习之外的类别特征. Zhou 等人^[15]应用对比学习来最大化属性上的互信息, 并通过随机屏蔽属性和项目实施序列增强. Qiu 等人^[16]利用基于 dropout 的模型级增广的对比正则化来构建对比样本. Cai 等人^[17]提出一种图对比学习范式, 利用奇异值分解进行对比增强. Liu 等人^[18]对数据增强算子进行改进, 并结合先进的顺序推荐模型实施推荐. Dang 等人^[19]在 Qiu 等人研究基础上考虑了用户会话序列中时间间隔对推荐结果的影响, 进一步改良了数据增强方法, 提出 TiCoSeRec 模型. 方阳等人^[20]运用异质信息网络搭建对比元学习的框架, 缓解冷启动问题. 钱忠胜等人^[21]提出了一种双分支协同增强的推荐模型, 动态调整用户的偏好和项目的热度. Xin 等人^[22]使用 3 种新的数据删除策略, 并根据会话相似度将会话划分为单独的子模型训练, 以提高学习效率. Zhong 等人^[23]设计多阶段分层去噪机制和全局关系编码器的运用, 这显著提高了序列数据的恢复质量.

数据增强的推荐模型为提高增强数据的质量, 结合时空信息、自监督学习和对比学习等技术, 设计多种增强策略和算子来变换用户会话来丰富信息, 缓解数据稀疏性问题, 但大多数对比学习框架形式单一, 限制了模型的泛化能力.

● 基于图神经网络的推荐方法. 该类模型先利用用户会话信息构建出序列图, 以此为基础使用图神经网络来对项目间的信息进行编码, 得到相应的嵌入表示; 再利用项目的嵌入表示计算得到会话的嵌入表示, 以此为最后推荐任务的基础. Wang 等人^[10,11]基于会话图利用 GNN 方法学习项目间的转换关系. Qiu 等人^[3]通过多头加权注意力图神经网络学习项目的嵌入表示, 可更好地捕捉会话中项目的最佳顺序. Wang 等人^[11]通过对单个会话图和全局会话图进行图卷积, 学习会话级和全局级嵌入. 随着深度学习的蓬勃发展, 超图神经网络受到广泛关注, 超图提

供了一种自然的方法来捕获复杂的高阶关系. Feng 等人^[24]以及 Yadati 等人^[5]最早将超图应用于 GCN. Jiang 等人^[6]开发出一种动态超图神经网络用于捕捉超图中的动态演化关系. Bandyopadhyay 等人^[7]开发了直线超图卷积网络, 保留了超图的全局连接性, 使模型能在超图结构上有效学习节点表示. Wang 等人^[25]提出的 HyperRec 使用超图建模用户的短期偏好, 但未利用超边间的信息, 也不是为基于会话的场景设计的. 而 Xia 等人^[8]在超图的基础上, 将会话关系建立成超边, 从而提取会话间的关系, 提出了 DHCN 模型. Zhao 等人^[26]将动态超图与聚类分布采样结合以进行多层次推荐. Yu 等人^[27]挖掘用户交互序列中重复项目的重要关系, 以重复项目作为超边关联用户交互序列. 虽然上述模型考虑了全局兴趣和顺序的因素, 但是仍未充分挖掘会话内的侧信息, 且对于会话间的信息挖掘不够深入.

基于图神经网络的推荐模型通过构建序列图, 图卷积等技术, 学习项目间的互动信息以预测用户兴趣. 研究者们提出了多种超图模型, 如 FGNN、GCE-GNN 和 HyperRec, 以捕捉会话内项目的最佳顺序和全局嵌入. 尽管这些模型考虑了顺序因素和全局兴趣, 但会话内时间间隔的侧信息挖掘仍不充分, 且在实际应用中忽略了会话间的隐藏高阶关系, 导致会话信息的多样性不足, 影响推荐结果的准确性.

综上所述可知, 多数模型并未关注到会话内时间间隔等侧信息的利用^[28], 从而造成对用户会话建模信息不足的现象, 影响推荐准确性; 同时, 现有的会话模型未很好地处理会话间的隐藏高阶信息; 此外, 大多使用对比模型的框架较单一, 使其泛化性不强. 针对上述问题, 本文提出一种结合时间间隔数据增强的对偶视图自监督会话推荐模型 TIDA-DSSR. 该模型深度挖掘会话内的时间间隔信息, 根据用户会话特征对其实施有针对性的数据增强, 更合理地建模用户兴趣偏好; 同时, 提出一种对偶视图编码器, 分别从会话间以及会话内处理会话信息, 结合时间间隔数据增强模块可更好地捕捉隐藏的高阶关系; 另外, 构建对比学习框架自适应地利用不同视图下的会话信息, 提升模型泛化性.

3 基于数据增强的对偶视图自监督会话推荐

为便于阐述, 对文中一些主要符号进行说明, 如表 1 所示.

表 1 主要符号及含义

符号	含义
$U = \{u_1, u_2, \dots, u_N\}$	用户集
$V = \{v_1, v_2, \dots, v_M\}$	项目集
$T_u = \langle t_1, t_2, \dots, t_m \rangle$	用户 u 会话项目间隔序列
p	数据增强项目插入位置
A	邻接矩阵
D	超图节点度矩阵
B	超图边度矩阵
S_u	用户 u 的原始会话集
S_u^a	用户 u 的增强会话集
s_u	用户 u 的原始会话序列
s_u^a	用户 u 的增强会话序列
$h_{s_u}^a$	增强会话 s_u^a 的超图表示
$f_{s_u}^a$	增强会话 s_u^a 的Transformer表示
f_{s_u}	原始会话 s_u 的Transformer表示

大多会话推荐模型重点关注顺序推荐^[29], 但未充分考虑会话内项目的时间间隔信息对会话推荐的影响. 同时, 在多个会话间还隐藏着许多对推荐或将起到重要作用的高阶关系, 且基于图神经网络的会话推荐模型能有效地挖掘这种隐藏高阶信息. 近来, 对比学习常用于融合不同视角的信息, 然而目前的对比学习应用框架形式单一, 使模型泛化性不足. 基于此, 提出一种自监督推荐模型利用时间间隔信息针对性地进行数据增强, 接着构建对偶视图编码器多角度编码会话嵌入, 最后通过对比学习融合多种会话信息. 图 1 是本文模型的总体框架, 其工作原理如下.

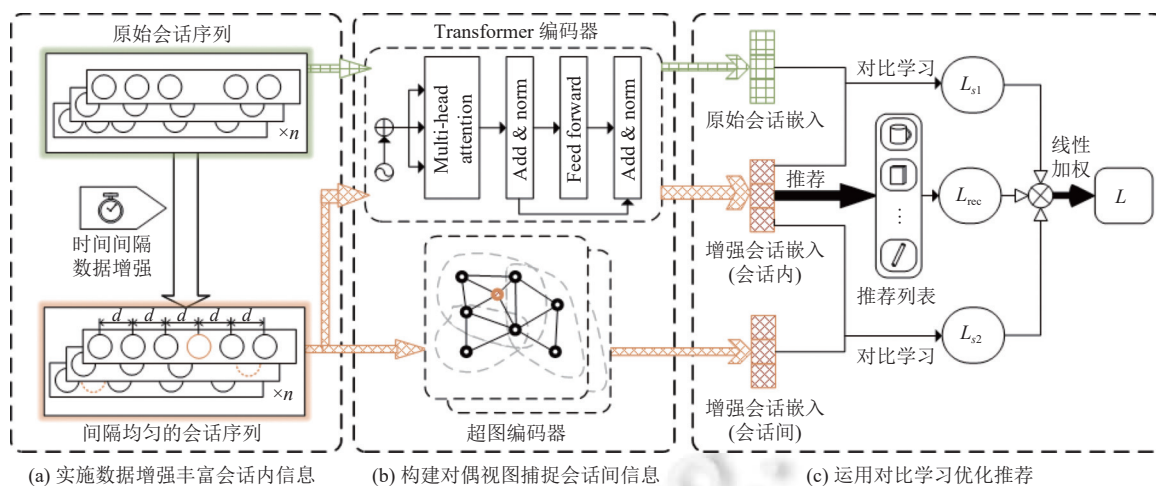


图1 结合时间间隔数据增强的对偶视图自监督会话推荐框架

1) 利用时间间隔信息实施数据增强,以丰富会话内信息.统计用户会话序列中项目间的时间间隔信息并计算出方差,以此为依据对不同会话进行不同方式的数据增强,使会话成为时间间隔均匀的会话序列,充分挖掘会话内的侧信息.

2) 构建对偶视图进行编码,以捕捉会话间信息.考虑用户会话内以及相似会话间的关系,利用超图卷积网络编码器与 Transformer 编码器组成的对偶视图编码器,从两种不同视角分别对原始会话以及间隔均匀的会话进行编码,以提取不同视角下会话的多样性信息.

3) 运用对比学习优化模型,并生成推荐列表.将原始会话与时间间隔数据增强后的会话输入对偶视图编码器得到不同视角下的会话嵌入,以数据增强后的会话嵌入(会话内)作为主体生成推荐列表,用其他视角下的会话嵌入(原始会话、会话间)作为对比信号加入自监督学习,以优化模型.

3.1 实施数据增强丰富会话内信息

数据增强一直是生成对比学习所需自监督信号的重要手段.有效的数据增强策略不仅可针对性地生成对比学习信号,也能挖掘会话内的侧信息,使会话嵌入信息更丰富,从而提高推荐准确性. TiCoSeRec^[19]首次将会话内项目间的时间间隔时长信息融入数据增强策略中,将原始会话增强为时间间隔均匀的会话,并充分证明了其能有效地提升推荐效果.本研究受此文启发,从会话内项目间的时间间隔时长信息入手对数据增强的4种方法(即 Ti-Insert 操作、Ti-Crop 操作、Ti-Mask 操作、Ti-Substitute 操作)实施改进.我们在 TiCoSeRec 模型的基础上有针对性地调整与改良,以提升其在匿名会话推荐场景中的适应性与泛化能力.首先,在 Ti-Insert 操作中,引入基于统计分布的动态插入比,同时结合项目间的相似度,确保插入项目符合上下文环境并增补潜在的语义信息.其次,针对 Ti-Crop 操作,通过动态调整裁剪比,限制对短会话的干扰,保持核心行为序列的完整性,同时增强生成后数据的多样性.此外,在 Ti-Mask 操作中,基于时间间隔和语义相关性引入优先级筛选机制,精准选择需屏蔽的项目,避免随机屏蔽破坏上下文信息的完整性,同时保留关键信息以辅助模型学习.而在 Ti-Substitute 操作中,不仅考虑项目间的相似性,而且还结合上下文语义的一致性,确保替换项目既丰富了语义信息又保持会话的整体逻辑.通过这些改进方法,进一步提升了增强数据的质量和多样性,为匿名会话推荐任务提供了更加高效的、鲁棒性更强的解决方案.

在数据增强过程中,本研究避免了使用 Reorder 操作,因为模型的超图卷积编码器需要依赖用户会话中项目顺序来识别不同会话间的隐性关联. Reorder 操作可能会破坏会话中项目顺序,导致模型在卷积过程中错误地连接无关的邻居节点,从而影响高阶隐藏关系的提取.排除 Reorder 操作有助于保持数据集的真实性和一致性,减少对原始数据结构的破坏,更准确地捕捉用户行为的时序特征,避免信息丢失和噪声增加.用户浏览习惯的时序规律对

于理解其兴趣变化至关重要. 例如, 商品的连续查看可能预示着购买决策, 而视频观看顺序则反映了用户的偏好. 保留这些时序信息使模型能够更精确地提取用户兴趣的动态变化, 实现更个性化的推荐.

假定有用户集为 U 、项目集为 V , 每个用户 $u \in U$ 按时间顺序交互的项目序列 $s_u = \langle v_1, \dots, v_j, \dots, v_N \rangle$, 其中 $s_u \in S_u$, $v \in V$, 其对应的时间间隔序列为 $T_u = \langle t_1, \dots, t_j, \dots, t_{N-1} \rangle$. 这一步可概括为, 对用户的交互项目序列 s_u , 根据其时间间隔 T_u 实施数据增强得到新的项目序列 s_u^a . 对话集 S_u 中的所有非均匀会话实施数据增强, 从而得到时间间隔均匀的会话集 S_u^a , 如公式 (1) 所示:

$$S_u^a = \text{DataAug}(T_u, S_u) \quad (1)$$

先根据时间间隔的标准方差对所有会话进行排序. 方差越小, 排名越前. 将排名最前的一定比例 (用 σ 表示) 的会话序列标记为均匀的, 其余为需进行数据增强的非均匀会话序列. 若共有 m 个会话序列, 则就有 $m(1-\sigma)$ 个序列为非均匀的. 为生成均匀的会话序列, 本文采用 4 种利用时间间隔时长信息对会话序列进行数据增强的操作.

1) Ti-Insert 操作. 先对序列 T_u 的时间区间从大到小排序, 保留前 $k = \beta N$ 个作为目标位置. 其中, N 为会话序列长度, $\beta \in [1, 0]$ 为插入比. 接着在每个目标位置 p 插入一个项目, 该项目尚未与用户交互, 但与序列中的上下文相关. 进一步地, 我们根据插入项目与相邻项目的相似度自适应地调整它们的时间间隔, 将插入项目与更相似项目的时间间隔设置得更小, 即, $t_p = t_i + \gamma(t_{i+1} - t_i)$, 其中 $\gamma \in [0, 1]$ 是插入比, 可通过观察时间间隔的统计分布动态调节, 模拟真实的交互行为.

例 2: 用户第 1 天购买了书、钢笔, 第 2 天购买了信纸. 我们就认为在此购物序列之间可插入一个相关项目“墨水”, 使购物序列均匀完整.

$$(p_1, p_2, \dots, p_k) = \text{top_k_indices}(\text{descend}(T_u)) \quad (2)$$

2) Ti-Crop 操作. 先选择一个截止位置 p , 在 p 处将长为 $c = \eta n$ 的连续序列截断为一个新的序列, 其中 $\eta \in [0, 1]$ 是一个比例系数, n 是会话序列长度. 多次操作生成一组长度均为 c 的子序列 (记为 (s_u, c)), 并计算每一个子序列的标准差. 然后选择标准差最小的子序列作为操作后的新序列.

例 3: 用户第 1 天购买了书、钢笔、信纸, 第 2 天购买了衬衫、领带. 我们就认为此购物序列需要分为两个子序列.

$$p = \underset{p \in [1, N-1]}{\text{argmin}} \text{std}(\text{sub_sequence}(s_u, c, p)) \quad (3)$$

对于短会话或稀疏会话 (项目数较少), 我们控制裁剪比 $c = \min\left(0.2, \frac{1}{\log(|s_u| + 1)}\right)$, 确保短会话裁剪比较小, 避免关键行为丢失.

3) Ti-Mask 操作. 与 Ti-Insert 操作类似, 都需要一个目标位置来进行变换会话内的项目. 但不同的是, Ti-Mask 操作是从原始序列中删除项目. 具体地, 该操作先对序列 T_u 的时间间隔时长从小到大排序, 然后保留前面的 n 个索引位置作为目标位置.

例 4: 用户第 1 天上午购买了书、钢笔、信纸, 下午购买了衬衫, 第 2 天购买了书签、书架. 我们就认为“衬衫”不属于整个购物序列, 选择删除该项目.

$$(p_1, p_2, \dots, p_n) = \text{top_h_indices}(\text{ascend}(T_u)) \quad (4)$$

在选取位置后, 根据与相邻项目的时间间隔和相似性计算出待删除项目的 Mask 优先级 $\text{pri_mask} = \frac{1}{1 + \text{sim}(v_i, v_{i+1}) \cdot (t_{i+1} - t_i)}$, 我们优先选取时间间隔大且相似度低的项目进行 Mask 操作, 避免会话中的重要交互项目被剔除.

4) Ti-Substitute 操作. 与 Ti-Mask 使用相同的策略来确定目标位置, 但不同的是, Ti-Substitute 操作要将位置上的项目替换成与上下文相关且用户未交互过的项目. 而且, 我们选取的替换项目与原项目具有相似的时间分布, 即, $|t_{\text{sub}} - t_{v_i}| \leq \kappa$, 确保替换后的项目符合原会话时序逻辑. 该操作能在对原会话实施改动最小的基础上, 获得新会话序列.

例 5: 用户购买了书、钢笔、信纸、书桌. 我们认为可把“书桌”替换为“书架”, 以谋求推荐对于阅读类偏好物品建立联系.

利用时间间隔信息进行数据增强的具体过程见算法 1.

算法 1. 基于时间间隔信息的数据增强算法.

输入: 用户集 U , 项目集 V , 用户会话序列 S_u , 项目表示;

输出: 数据增强后得到的时间间隔均匀的会话序列.

Begin

1. 初始化用户会话序列 S_u ; //用户交互会话
 2. **For** $s_u \in S_u$:
 3. 计算用户会话序列 s_u 中项目的时间间隔 $T_u = \langle t_1, \dots, t_j, \dots, t_{N-1} \rangle$;
 4. 计算会话时间间隔的方差 δ^2 ;
 5. **End For**
 6. 根据方差 δ^2 大小, 将用户会话序列排序 $S_u = \langle s_{u1}, s_{u2}, \dots, s_{uN} \rangle$;
 7. 根据方差大小对会话进行对应的数据增强 S_u^a ; //见公式 (2)–公式 (4)
 8. **Return** S_u^a ; //返回时间间隔均匀的会话序列
 9. **End**
-

在算法 1 中, 第 1 行得到用户的会话序列 S_u ; 第 2–5 行计算会话内项目的时间间隔 T_u , 再计算会话内项目时间间隔的方差 δ^2 ; 第 6 行根据方差 δ^2 的大小, 对所有的用户会话进行排序, 根据排序先后会选用不同的数据增强方法; 第 7 行对用户会话进行数据增强; 第 8 行得到数据增强后时间间隔均匀的用户会话.

3.2 构造对偶视图捕捉会话间信息

在推荐领域, 对比学习作为一种有效的自监督学习方法, 近年来受到广泛关注. Yu 等人^[30]的研究通过在嵌入空间中添加均匀噪声来创建对比视图, 以此进行对比学习, 取得显著的推荐效果提升. 受此启发, 我们进一步探索了如何利用时间间隔时长信息来改进数据增强策略, 并在此基础上构建对偶视图以挖掘会话间的隐藏高阶关系.

所提模型首先对原始会话数据进行基于时间间隔的数据增强处理, 包括 Ti-Insert、Ti-Crop、Ti-Mask 和 Ti-Substitute 这 4 种操作. 这些操作旨在模拟用户在不同场景下的行为模式, 从而生成多样化的增强会话序列. 随后, 我们构建对偶视图编码器捕捉会话间的隐藏高阶关系. 先将这些增强后的会话序列构成一个超图, 其中每个节点代表一个项目, 边则表示项目之间的时序关系或相似度. 这样的超图结构不仅保留了原始会话中的时间间隔信息, 而且还引入了额外的上下文信息, 有助于模型更好地理解用户的兴趣变化过程. 为确保对比学习的有效性, 我们还采用 Transformer 编码器同时处理数据增强后的会话以及原始会话. Transformer 编码器的强大特征提取能力使得它可从复杂的序列中捕捉到用户的兴趣变化, 为超图视图提供高质量的对比信号.

3.2.1 超图卷积网络编码

所提模型 TIDA-DSSR 结合了时间间隔数据增强模块和对偶视图编码模块, 能够有效地捕捉会话间的隐藏高阶关系. 通过第 3.1 节的数据增强过程, 会话信息与相似的项目已建立了关联, 这为深入挖掘会话间的隐藏高阶关系提供了基础.

超图结构在 TIDA-DSSR 模型中至关重要, 其超边能连接多个节点, 与会话中多个项目的天然特性相匹配. 这种结构优化了会话嵌入的建模, 且可有效地表示复杂的多对多关系. 具体来说, 超图结构不仅捕捉单个项目之间的直接联系, 还考虑项目间复杂的间接关系, 以更全面地理解用户行为. 同时, 超图结构还可缓解传统模型的长距离依赖问题, 通过超边连接增强了模型对长会话序列关联的捕捉能力. 故我们参照 Xia 等人^[8]的方法, 采用超图 $G_h = (V_h, E_h)$ 来表示所有会话, 捕捉会话中的这种关系. 形式上, 将会话表示为超边, 记为 $\{v_{s,1}, v_{s,2}, v_{s,3}, \dots, v_{s,m}\} \in E_h$, 且 $v_{s,m} \in V_h$. 在将会话信息转换为超图后, 单个会话内的所有项目均被连接, 这就对多对多的高阶关系进行了具体

化. 同时, 通过超图节点与超边的连接, 实现会话及其相似会话的关系提取. 与 Xia 等人^[8]方法不同的是, 本文不是采用超边图卷积网络提取高阶关系, 因为这会融合过多的邻居信息, 从而易使用户的兴趣产生漂移, 所以我们采用 Transformer 编码器提取原始会话中更重要的顺序信息 (见第 3.2.2 节).

在超图构建之后, 使用超图卷积网络来捕捉会话级的隐藏高阶关系. 定义超图上的卷积操作的主要挑战是如何传播项目的嵌入. 我们使用的光谱超图卷积来构建超图卷积网络, 如公式 (5) 所示:

$$e_{s_u,i}^{(l+1)} = \sum_{i=1}^N \sum_{e=1}^M A_{ie} A_{je} W_{ee} e_{s_u,i}^{(l)} \quad (5)$$

其中, $e_{s_u,i}^{(l+1)}$ 代表在超图卷积 l 层会话 s_u 中项目 v_i 的嵌入表示, A 为邻接矩阵, W 为权重矩阵 (它的每个超边赋予相同的权重 1), M 为会话中的项目数, N 为相关联的会话数. 超图卷积网络的矩阵带行归一化的表达式, 如公式 (6) 所示:

$$H_{s_u,i}^{(l+1)} = D^{-1} A W B^{-1} A^T H_{s_u,i}^{(l)} \quad (6)$$

其中, $H_{s_u,i}^{(l+1)}$ 代表在超图卷积 l 层会话 s_u 中的项目 v_i 嵌入矩阵, B (或 D) 为边 (或节点) 度矩阵, A 为邻接矩阵, W 为权重矩阵. 超图卷积可看作是对超图结构进行“节点-超边-节点”的两阶段特征变换, 即“项目-会话-项目”的信息传递. 节点到超边的信息聚合由节点矩阵与邻接矩阵的乘法 $A^T H^{(0)}$ 完成, 后再乘以邻接矩阵 A , 这样将信息从超边传递到节点. 将 $H^{(0)}$ 通过 L 层超图卷积后, 对每一层得到的项目嵌入实施平均化操作, 得到最终的项目嵌入 $H_{s_u,i} = 1/(L+1) \sum_{l=0}^L H_{s_u,i}^{(l)}$.

在空间上传递信息时, 时间信息是不可忽视的. 而在会话序列中, 时间信息往往体现为项目的位置前后关系. Vaswani 等人^[31]认为位置嵌入是 Transformer 中引入的一种有效技术, 在许多领域被应用于项目位置信息的存储与运用. 参考 Xia 等人^[8]的方法, 我们设计一个可学习的位置矩阵 $Pos_r = [pos_1, pos_2, pos_3, \dots, pos_m]$, 其中 m 是当前会话长度, 通过该矩阵与会话嵌入拼接, 使时间信息融入会话嵌入. 在增强会话 $s_u^a = \langle v_1, v_2, v_3, \dots, v_m \rangle$ 中, $e_{s_u,i}^*$ 为增强会话 s_u^a 中 v_i 的嵌入表示, 如公式 (7) 所示:

$$e_{s_u,i}^* = \tanh(W_1 [e_{s_u,i} \parallel pos_{m-i+1}] + b) \quad (7)$$

其中, $W_1 \in \mathbb{R}^{d \times 2d}$, $b \in \mathbb{R}^d$ 是可学习的参数. 通过对会话包含的所有项目嵌入进行平均操作来表示此会话的嵌入, 以细化增强会话 s_u^a 的嵌入表示, 如公式 (8) 所示:

$$h_{s_u}^a = \frac{1}{m} \sum_{i=1}^m e_{s_u,i}^* \quad (8)$$

其中, $h_{s_u}^a$ 是增强会话 s_u^a 的嵌入, $e_{s_u,i}^*$ 是增强会话 s_u^a 中第 i 个项目的嵌入.

上面经过超图卷积网络编码器获得的会话嵌入包含了会话间的关系, 然而, 准确预测用户兴趣还需提取会话内信息.

3.2.2 Transformer 编码

Transformer 编码器具备卓越的特征提取能力, 能够从复杂的序列数据中识别并捕捉到用户的兴趣变化. 这种能力使得它可为超图视图提供高质量的对比信号, 以增强模型的鲁棒性. 利用 Transformer 编码器可有效地筛选出会话中与目标用户兴趣不相关的信息, 不仅能够更全面地提取会话内的短期兴趣, 还能在捕捉会话间的潜在高阶关联上取得更大突破. 这种全局感知与高效计算的结合, 使其在复杂的推荐任务中表现优异, 为推荐任务提供了强有力的技术支持. 本文使用 Transformer 编码器提取会话内信息, 将 Transformer 编码器表示为 $SeqEnc(\cdot)$, 如公式 (9):

$$f_{s_u}, f_{s_u}^a = SeqEnc(s_u, s_u^a) \quad (9)$$

其中, f_{s_u} 表示原始会话嵌入, $f_{s_u}^a$ 表示增强会话嵌入, s_u 是原始会话, s_u^a 是增强会话. Transformer 编码器由一个序列块和一个全连接前馈网络构成, 每个编码器层包括多头自注意力层 (multi-head self-attention layer) 与前馈全连接层 (feed forward layer) 两个子层. 多头自注意力子层将会话序列作为输入计算 query、key 和 value, 通过缩放点积操作实现注意力计算, 并采用多套线性变换映射矩阵提取会话序列的不同表示. 前馈全连接子层包含两个线性变换和一个 ReLU 激活函数, 以对会话序列建模. 这两个子层均使用了残差连接和层归一化, 以提升性能和加速收敛.

过程. 即, 对于输入的会话序列, 首先使用 3 组可学习的权重矩阵产生 query、key 和 value 表示; 然后计算 attention 分数, 进行 Softmax 操作生成权重系数; 最后加权求和 value 值得到最终的 self-attention 表示. 这一过程使模型可提取会话序列内项目的位置信息. 前馈全连接子层可捕获序列的局部特征, 与多头自注意力子层的输出表示互补. 这样, 将增强会话和原始会话通过 Transformer 编码器获得相应嵌入表示以实现下游任务.

本节阐述了从会话间 (见第 3.2.1 节) 及会话内 (见第 3.2.2 节) 两个不同视角构建对偶视图编码器, 对原始会话和增强会话实施多角度建模. 对偶视图编码器构建的具体过程见算法 2.

算法 2. 对偶视图编码器构建.

输入: 用户会话序列 S_u , 数据增强后的会话序列 S_u^a , 超图卷积层数 θ ;

输出: 不同角度下的会话嵌入.

Begin

1. 初始化超图 G_h , 包括点集合 V_h , 边集合 E_h ;
 2. 构建邻接矩阵 A , 点的度矩阵 D , 边的度矩阵 B ;
 3. **For** $\vartheta \in \theta$:
 4. 计算每层项目嵌入 $e_{s_u,i}^{(l+1)}$; //见公式 (5)
 5. 获取项目嵌入归一化 $H_{s_u,i}^{(l+1)}$; //见公式 (6)
 6. **End For**
 7. 平均化计算项目嵌入 $H_{s_u,i}$;
 8. 嵌入位置信息得到最终项目嵌入 $e_{s_u,i}^*$; //见公式 (7)
 9. 将会话中的项目嵌入平均计算获得增强会话嵌入 $h_{s_u}^a$; //见公式 (8)
 10. 计算 Transformer 编码器下的增强会话嵌入, 原始会话嵌入 $f_{s_u}, f_{s_u}^a$; //见公式 (9)
 11. **Return** $h_{s_u}^a, f_{s_u}, f_{s_u}^a$; //返回 3 种会话嵌入, 见公式 (8)、公式 (9)
 12. **End**
-

在算法 2 中, 第 1 行初始化超图 G_h , 其中包括点集 V_h 和边集合 E_h ; 第 2 行根据超图 G_h 计算出邻接矩阵 A , 以及点的度矩阵 D 和边的度矩阵 B ; 第 3–6 行计算超图卷积网络每层中的项目嵌入 $H_{s_u,i}^{(l+1)}$; 第 7 行将每层的项目嵌入平均化得到项目嵌入 $H_{s_u,i}$; 第 8 行表示在会话嵌入中添加位置信息得到最终项目嵌入 $e_{s_u,i}^*$; 第 9 行将增强会话中的项目嵌入平均计算获得其嵌入 $h_{s_u}^a$; 第 10 行计算 Transformer 编码器中的会话嵌入 f_{s_u} 和 $f_{s_u}^a$.

3.3 运用对比学习优化推荐

将原始会话信息与增强后的会话信息送入对偶视图编码器, 可生成包含不同信息特征的 3 种会话嵌入, 即原始会话嵌入、会话内的增强会话嵌入、会话间的增强会话嵌入. 这 3 种嵌入均代表同一会话, 但从不同的视角提取了关键特征, 其中包括会话内的短期行为模式和会话间的高阶关联. 3 种嵌入彼此之间存在一对一的映射关系, 这种关系使得每轮迭代中可将相同会话的不同嵌入视为正样本, 而随机抽取的其他会话嵌入作为负样本. 这样, 通过对比学习融合多视角的信息, 实现对会话嵌入的深度优化.

该操作的优势在于, 高效融合了会话内与会话间信息. 从会话内角度看, 通过编码会话序列中时序信息和行为模式, 捕捉用户短期兴趣的动态演化, 更精准地预测用户行为; 而从会话间角度看, 通过利用时间间隔信息数据增强和超图结构, 挖掘会话间潜在的高阶关系, 以揭示用户行为的全局上下文关联. 而且, 所提模型 TIDA-DSSR 通过对偶视图的编码方式, 生成多视角嵌入, 从多维度表达会话特征, 并在对比学习中互相强化, 弥补了单一视图编码可能遗漏的重要信息.

通过对偶视图编码器产生高质量的正负样本并进行多视角对比学习, TIDA-DSSR 模型不仅能强化同一会话中不同视角嵌入之间的语义一致性, 还可有效区分不同会话的特征, 避免信息干扰. 这种设计既保留了会话内短期

兴趣的细粒度刻画能力,又充分挖掘了会话间的潜在高阶关联,使模型在兼顾实时性和全局性的同时提升泛化能力和推荐性能,为复杂的推荐任务提供更全面的支持.

对于会话间信息的融合,本文将标准二元交叉熵损失函数作为学习目标,可直观地衡量样本对的相似性或差异性,以便更有效地进行正负样本的区分,提升模型的优化效果,如公式 (10) 所示:

$$L_{s1} = -\log(\delta(f_{s_u}^a, h_{s_u}^a)) - \log(1 - \delta(f_{s_u}^a, h_{s_c}^a)) \quad (10)$$

其中, $f_{s_u}^a$ 是 Transformer 编码器下的增强会话嵌入 (锚点), $h_{s_u}^a$ 是超图编码器下的对应增强会话嵌入 (正样本), $h_{s_c}^a$ 是超图编码器下的其他增强会话嵌入 (负样本), δ 为非线性激活函数.

然而,仅使用经过数据增强后的会话特征进行推荐任务,会导致对比学习框架的形式单一,使得模型在早期训练时会受到项目建模不准确的影响,在捕捉用户偏好变化上存在局限性,从而限制模型的泛化能力.因此,本文使用原始会话嵌入与增强会话嵌入进行对比学习,以避免无关会话项目的影响,提升模型的泛化能力.对于原始会话信息,这里采用 NT-Xent 损失函数来提取,有助于提升模型的特征表示能力,使得学习到的特征在下游任务中表现更佳,如公式 (11) 所示:

$$L_{s2} = -\log \frac{\exp(\text{sim}(f_{s_u}^a, f_{s_u}^a))}{\sum_{c \neq u} \exp(\text{sim}(f_{s_u}^a, f_{s_c}^a))} \quad (11)$$

其中, $\text{sim}(\cdot)$ 是两个会话嵌入间相似性度量的点积操作, $f_{s_u}^a$ 是对应原始会话嵌入 (正样本), $f_{s_c}^a$ 是其他原始会话的嵌入 (负样本).

由于经过数据增强的会话中包含时间间隔时长信息,且来自 Transformer 编码器的输出嵌入中带有位置信息,因此经过 Transformer 编码器的增强会话嵌入会融合多角度会话内项目的关系,我们将其作为推荐任务的主体.这里采用对数似然损失函数进行推荐,对编码器进行优化,如公式 (12) 所示:

$$L_{\text{rec}}(u, t) = -\log(\delta(f_{s_u}^a, e_{s_{u,t+1}})) - \sum_{c \neq u} \log(1 - \delta(f_{s_u}^a, e_{s_{c,i}})) \quad (12)$$

其中, $L_{\text{rec}}(u, t)$ 为会话 s_u 中预测第 t 位项目的损失函数, δ 为非线性激活函数, $e_{s_{u,t+1}}$ 为会话中下一项目 v_{t+1} 的嵌入, $e_{s_{c,i}}$ 不是相应会话 s_u 中的项目嵌入.项目的嵌入 e 从 SeqEnc 嵌入表中检索,该表与 Transformer 共同优化.

在推荐任务和对比学习任务中获得的损失函数均可对模型进行优化.为通过对比学习提高推荐性能,我们利用了一种多任务策略来联合优化模型,如公式 (13) 所示:

$$L = L_{\text{rec}} + \lambda_1 L_{s1} + \lambda_2 L_{s2} \quad (13)$$

其中, λ_1 、 λ_2 是超参数,用于控制对比学习任务的强度.自监督会话推荐的具体过程见算法 3.

算法 3. 结合时间间隔数据增强的对偶视图自监督会话推荐算法.

输入: 用户数据增强后的会话信息集 S_u^a , 对偶视图得到的会话嵌入 $h_{s_u}^a, f_{s_u}^a, f_{s_u}^a$;

输出: 会话推荐列表 $List$.

Begin

1. 获取数据增强后的会话 S_u^a ; // 见算法 1
 2. **For** 会话 $s_u^a \in S_u^a$:
 3. 获取对偶视图编码会话的不同嵌入 $h_{s_u}^a, f_{s_u}^a, f_{s_u}^a$; // 见算法 2
 4. 在不同视图间进行对比学习, 获得损失函数 L_{s1}, L_{s2} ; // 见公式 (10), 公式 (11)
 5. 实施推荐任务获得推荐列表 $List_{s_u}$ 及损失函数 L_{rec} ; // 见公式 (12)
 6. 计算联合优化损失函数 L , 并更新模型参数; // 见公式 (13)
 7. **End For**
 8. 获取推荐列表 $List \leftarrow List_{s_u}$; // 模型预测
-

9. **Return** 会话推荐列表 $List$;

10. **End**

在算法 3 中, 第 1 行获得数组增强后的会话集 S_u^a . 第 2-3 行通过对偶视图编码器获得增强会话 s_u^a 的不同会话嵌入 $h_{s_u^a}^a, f_{s_u^a}, f_{s_u^a}^a$; 第 4 行通过不同视图间的对比学习获得不同的损失函数 L_{s1}, L_{s2} ; 第 5 行通过进行推荐任务获得当前会话的推荐列表 $List_{s_u}$ 与损失函数 L_{rec} ; 第 6-7 行结合 L_{s1}, L_{s2} 与 L_{rec} 联合计算损失函数 L , 并更新模型参数; 第 8-9 行获取推荐列表 $List$ 并输出.

4 实验及其分析

为验证 TIDA-DSSR 的有效性, 本文选择 4 个数据集进行综合实验对比分析, 重点回答下面几个问题.

RQ1: 与经典的、最新的模型相比, 本文模型的推荐效果如何?

针对此问题, 在第 4.4.1 节设置对比实验, 分别将本文模型与 11 种经典模型就性能方面在 4 个数据集下做对比, 阐述本文模型的优势.

RQ2: 模型的各个构件是否有存在的必要性?

针对此问题, 在第 4.4.2 节设置消融实验, 就基于时间信息的数据增强模块、超图卷积模块、原始会话对比学习模块这 3 个主要构件分别分析其对总体模型的作用.

RQ3: 本文模型的主要构件用到其他模型中是否有效?

针对此问题, 在第 4.4.3 节设置构件组合合理性实验, 分析本文模型的主要构件用于其他不同对比模型中对这些模型性能的影响, 以阐明所提模型的构件组合是合理的.

RQ4: 主要的超参数对于模型效果有何影响?

针对此问题, 在第 4.4.4 节设置参数敏感度实验, 分析主要超参数在不同数据集下对模型性能的影响, 以进一步优化模型.

RQ5: 模型的时间复杂度如何?

针对此问题, 在第 4.4.5 节设置模型复杂度分析, 计算本文模型的时间复杂度, 并于主流模型比较, 阐述本文模型复杂度情况.

RQ6: 模型框架的扩展性如何?

针对此问题, 在第 4.4.6 节设置模型框架分析, 就对比学习框架在复杂的、多样化的推荐场景时的扩展性进行了详细探讨.

4.1 数据集与环境设置

本文选取的 4 个公开数据集均来自于真实业务场景, 分别是 Beauty, Sports, Home, Yelp. 这 4 个数据集覆盖了各种用户需求和偏好, 例如美食、美容、运动和家居等. 这意味着推荐系统可针对不同用户提供更具个性化的推荐. 它们均为近几年被大量学者使用和研究的经典数据集, 统计结果如表 2 所示.

表 2 数据集信息

统计数据	Beauty	Sports	Home	Yelp
用户数	22 363	35 598	66 370	19 855
项目数	12 101	18 357	28 237	14 541
交互次数	198 502	296 337	551 682	207 045
稀疏性 (%)	99.92	99.95	99.97	99.92
平均会话长度	8.87	8.32	8.29	10.42

1) Beauty (<https://github.com/KingGugu/TiCoSeRec/tree/main/data>): 化妆品购买数据. 其收集了某化妆品购物网站的真实用户行为日志, 记录了用户的浏览、加购、订单等行为序列, 以及网站商品的相关信息.

2) Sports (<https://github.com/KingGugu/TiCoSeRec/tree/main/data>): 体育用品购买数据. 其包含某体育用品在线购物网站的用户点击流和交易数据, 并记录了用户浏览网站、点击商品以及订单的时间序列.

3) Home (<https://github.com/KingGugu/TiCoSeRec/tree/main/data>): 家居购物数据. 其记录了用户在家居商品网站的浏览、点赞、加购、订单等交互行为序列.

4) Yelp (<https://www.yelp.com/dataset>): 餐厅点评数据. 数据标签包括用户对商家的评论及评分等.

所提模型 TIDA-DSSR 采用 PyTorch 框架, Windows 10 64 位操作系统, PyCharm 2022, Python 3.7, torch 1.11.0, numpy 1.20.1, 内存 64 GB, CPU 为 12th Gen Intel(R) Core(TM) i7-12650H 2.30 GHz, GPU 为 NVIDIA GeForce RTX 4060. 参照文献 [19], 本文模型的嵌入向量维度设置为 128, 训练批次大小为 512, 学习率为 0.001, 使用 Adam 作为优化器优化模型, 最大序列长度设置为 50, Softmax 温度值为 1.0, 权重衰减参数为 $1E-6$, epoch 为 300, 并采取早期停止策略, 即若连续 150 个周期内验证数据集上的 $HR@20$ 和 $NDCG@20$ 指标未增加, 则提前终止训练.

4.2 评价指标

在评估会话推荐系统性能时, 我们采用 HR (hit ratio) 和 $NDCG$ (normalized discounted cumulative gain) 作为主要评价指标. HR 的选取源于其直观反映推荐准确性的能力, 直接衡量用户感兴趣的项目是否被成功推荐. 而 $NDCG$ 则因其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, 更能体现推荐系统的综合性能. 这两个指标能反映对推荐系统效能的全面评价.

1) HR , 命中率, 主要强调模型推荐的准确性, 如公式 (14) 所示:

$$HR = \frac{M}{N} \times 100\% \quad (14)$$

其中, M 为推荐任务命中的项目数, N 为推荐任务给出的总项目数.

2) $NDCG$, 归一化折损累计增益, 是一种衡量推荐结果质量的关键指标, 如公式 (15) 所示:

$$NDCG = \sum_{u \in U_{\text{test}}} \frac{1}{Y_u} \sum_{i=1}^K \frac{2^{t_i-1}}{\log_2(i-1)} \quad (15)$$

其中, U_{test} 为测试用户集; Y_u 为用户 u 的最大 $NDCG$ 值; K 为已推荐项目数; $t_i = 1$ 表示击中, $t_i = 0$ 表示未击中.

4.3 对比模型

下面将所提模型 TIDA-DSSR 与相关模型 LightGCN^[32], STAMP^[33], BERT4Rec^[34], DHCN^[8], CL4SRec^[9], CoSeRec^[18], DuoRec^[16], TiCoSeRec^[19], SSDRec^[23], S⁴Rec^[35], SparseEnNet^[36] 做对比实验, 阐明本文模型的优势. 这几种作为对比的深度模型采用了不同方法, 包括基于图神经网络的方法、基于强化学习的方法、基于自注意力机制的序列建模方法、基于自监督学习的方法等.

1) LightGCN: 为 GCN 推荐方面的代表性方法之一, 主要特点是移除了 GCN 中的特征变换和聚合函数组件, 使设计更简单、高效.

2) STAMP: 是将用户短期与长期兴趣融为一体的注意力机制运用于推荐系统的典型实例, 能根据用户的短期与长期兴趣相似性生成个性化的推荐.

3) BERT4Rec: 是将 BERT 运用于推荐系统的开创性研究, 通过使用双向编码器来建模用户行为序列中的上下文信息和长期依赖关系.

4) DHCN: 一种基于深度学习的会话推荐模型, 它利用动态超图结构来捕捉用户与项目之间的复杂关系, 并结合自监督学习提高推荐效果.

5) CL4SRec: 一种基于强化学习的会话推荐框架, 在会话中持续强化模型以建模用户动态的短期兴趣和漂移的长期偏好.

6) CoSeRec: 一种融合序列和语义信息的会话推荐模型, 通过协同学习同时建模用户的历史行为序列和项目内容的语义特征.

7) DuoRec: 一种双塔结构的推荐模型, 通过用户塔和项目塔分别学习用户和项目的低维表示, 并基于两者的

交互预测用户偏好.

8) TiCoSeRec: 一种基于自监督学习会话推荐模型, 目标是同时建模用户交互序列中的时间、上下文和顺序模式. 先构建一个时间感知的 GCN 来学习会话中的上下文和时间信息, 再利用一个基于自注意力机制的序列推荐网络来学习用户交互的顺序依赖关系, 从而综合序列、上下文和时间多个维度对用户偏好建模.

9) SSDRec: 一种最新的序列数据恢复自监督框架, 通过构建多关系图和分层去噪策略来高效地处理损坏或不完整的序列数据, 从而提高推荐结果的准确性.

10) S⁴Rec: 一种最新的对抗学习模型, 通过在线聚类和对抗学习来学习用户潜在意图, 并利用自蒸馏将有丰富行为数据的用户 (老师) 知识传递给行为数据稀疏的用户 (学生), 从而提高推荐性能.

11) SparseEnNet: 一种最新的鲁棒对抗生成方法, 旨在通过生成适应性更强的增强项目来缓解序列推荐中的数据稀疏问题.

上面 11 种对比模型与本文模型的详细特征描述如表 3 所示, 其中, “×”表示相应模型不存在该组件; “○”表示相应模型存在该组件.

表 3 各模型对比

模型名称	时间间隔信息	多视角	自监督学习
LightGCN	×	×	×
STAMP	×	×	×
BERT4Rec	×	○	×
DHCN	×	○	○
CL4SRec	×	×	○
CoSeRec	×	×	○
DuoRec	×	○	○
TiCoSeRec	○	×	○
SSDRec	×	×	○
S ⁴ Rec	×	×	○
SparseEnNet	×	×	○
本文模型	○	○	○

4.4 实验结果与分析

下面通过各节的实验结果来验证本文模型的优势及其有效性. 其中, 在第 4.4.1 节将本文模型与 11 种经典的、最新的相关模型进行实验对比, 详细分析并阐明本文模型的优势, 以回答 RQ1; 在第 4.4.2 节针对本文模型的 3 大构件设计了 5 种变体模型, 进行消融实验对比, 以回答 RQ2; 在第 4.4.3 节将本文模型的主要构件加入或替换到其他对比模型中, 观察这些模型的效果变化, 以回答 RQ3; 在第 4.4.4 节对模型的 4 个主要关键参数进行调优实验, 以回答 RQ4; 在第 4.4.5 节分析了各模型的复杂度, 以回答 RQ5; 在第 4.4.6 节就本文模型对比框架的可扩展性问题展开了讨论, 以回答 RQ6; 在第 4.4.7 节给出了案例分析, 直观展示了所提模型的效果.

4.4.1 对比实验 (RQ1)

在 4 个不同的经典数据集 (Beauty, Sports, Home 及 Yelp) 上进行实验, 保留用户与项目交互次数不小于 5 的会话信息, 并将每个用户的倒数第 1、第 2 个交互项目记录分别作为测试和验证数据, 其余的作为训练数据.

为评估所提模型 TIDA-DSSR 的效果, 各模型在不同数据集下的 HR 值和 $NDCG$ 值在 @10 和 @20 (在表中分别简写为 H, N) 的对比情况, 如表 4 所示.

1) 与未利用自监督学习的模型相比

① LightGCN 在推荐系统领域是非常重要的模型, 利用 GCN 学习用户和项目的表示, 使模型轻便化的同时, 也兼顾了不错的效果. 相较于 LightGCN 模型, 本文模型 TIDA-DSSR 在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 23.06%, 16.46%.

② STAMP 较 LightGCN 更加侧重于会话内的顺序, 使用 RNN 或 Transformer 等序列模型处理会话信息. 相

较于 STAMP 模型, 本文模型 TIDA-DSSR 在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 31.40%, 57.34%.

③ BERT4Rec 使用双向自注意力机制来建模用户的行为序列, 相比单向结构而言, 它能获得更准确的用户行为序列表示, 从而提高推荐性能. 相较于 BERT4Rec 模型, 本文模型 TIDA-DSSR 在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 31.30%, 34.09%.

表 4 各模型对比结果

数据集	指标	无自监督模型			自监督模型								性能提升 (%)	
		Light-GCN	STAMP	BERT-4Rec	DHCN	CL4SRec	DuoRec	CoSeRec	TiCo-SeRec	SSDRec	S ⁴ Rec	Sparse-EnNet		TIDA-DSSR
Beauty	H@10	0.0540	0.0508	0.0526	0.0564	0.0569	0.0686	0.0675	0.0737	0.0517	0.0719	<u>0.0762</u>	0.0842	10.50
	H@20	0.0803	0.0743	0.0825	0.0899	0.0885	0.1022	0.1015	0.1079	0.0776	0.1071	<u>0.1103</u>	0.1249	13.24
	N@10	0.0292	0.0283	0.0263	0.0316	0.0277	0.0359	0.0381	0.0421	0.0241	<u>0.0426</u>	0.0424	0.0468	9.86
	N@20	0.0358	0.0342	0.0338	0.0420	0.0356	0.0470	0.0467	0.0499	0.0306	0.0505	<u>0.0512</u>	0.0568	10.94
Sports	H@10	0.0378	0.0327	0.0332	0.0368	0.0414	0.0392	0.0421	0.0507	0.0259	0.0482	<u>0.0527</u>	0.0553	4.93
	H@20	0.0578	0.0474	0.0538	0.0532	0.0637	0.0630	0.0629	0.0767	0.0391	0.0656	<u>0.0779</u>	0.0822	5.52
	N@10	0.0200	0.0185	0.0164	0.0197	0.0215	0.0195	0.0241	0.0282	0.0122	0.0257	<u>0.0295</u>	0.0305	3.39
	N@20	0.0251	0.0222	0.0216	0.0264	0.0271	0.0256	0.0292	0.0345	0.0156	0.0292	<u>0.0351</u>	0.0373	6.27
Home	H@10	0.0096	0.0178	0.0165	0.0201	0.0184	0.0251	0.0232	<u>0.0267</u>	0.0183	—	—	0.0281	5.24
	H@20	0.0155	0.0247	0.0224	0.0322	0.0285	0.0347	0.0336	<u>0.0391</u>	0.0264	—	—	0.0412	3.92
	N@10	0.0051	0.0121	0.0089	0.0126	0.0092	0.0124	0.0136	<u>0.0153</u>	0.0118	—	—	0.0159	5.37
	N@20	0.0065	0.0133	0.0103	0.0158	0.0118	0.0152	0.0165	<u>0.0184</u>	0.0087	—	—	0.0196	6.52
Yelp	H@10	0.0542	0.0412	0.0508	0.0572	0.0581	0.0627	0.0613	<u>0.0658</u>	0.0479	—	0.0414	0.0667	1.37
	H@20	0.076	0.0668	0.0786	0.0836	0.0896	0.0976	0.0965	<u>0.1032</u>	0.0656	—	0.0678	0.1051	2.96
	N@10	0.0328	0.0214	0.0279	0.0294	0.0312	0.0345	0.0322	<u>0.0371</u>	0.0309	—	0.0209	0.0382	1.94
	N@20	0.0383	0.0278	0.0352	0.0392	0.0390	0.0432	0.0425	<u>0.0459</u>	0.0354	—	0.0275	0.0472	2.83

注: 本文模型TIDA-DSSR的实验数据以加粗字体标示, 为便于比较, 利用下划线来突显对比模型中表现最佳的数据, 最后一列给出本文模型相对于某一最佳对比模型的性能提升情况(粗体标示)

3 大经典模型 LightGCN、STAMP 和 BERT4Rec 分别从高阶协同信息、会话内顺序信息和用户行为上下文的角度考虑对项目建模, 但未深层次地挖掘用户会话序列内的时间信息, 且推荐任务本身存在的稀疏性问题也未得到改善. 因此, 本文模型 TIDA-DSSR 采用基于时间间隔的数据增强和对比学习相结合的方式应对上述问题.

2) 与自监督学习模型相比

① DHCN 使用超图来表示用户和项目之间的复杂交互信息, 可更好地捕捉到其中的高阶关联关系. 相较于 DHCN 模型, 本文模型在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 16.61%, 20.41%.

② CL4SRec 为缓解数据稀疏性问题第 1 次将对比学习运用到顺序推荐. 相较于 CL4SRec 模型, 本文模型在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 14.80%, 14.80%.

③ DuoRec 在模型中构建 Dropout 的模型级数据增强方法, 同时利用相似序列挖掘自监督信号. 相较于 DuoRec 模型, 本文模型在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 6.38%, 9.26%.

④ CoSeRec 在 DuoRec 基础上引入更多的数据增强操作. 相较于 CoSeRec 模型, 本文模型在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 8.81%, 11.06%.

⑤ TiCoSeRec 在 CoSeRec 的基础上对时间间隔时长分布进行了细致研究, 提出 5 个数据增强算子, 将非均匀序列转换为均匀序列, 从而更好地捕捉时间信息以提升推荐性能. 相较于 TiCoSeRec 模型, 本文模型在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 1.37%, 2.83%.

⑥ SSDRec 通过设计多关系图和分层去噪策略来高效地处理损坏或不完整的序列数据. 相较于 SSDRec 模型, 本文模型在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 53.55%, 34.75%.

⑦ S⁴Rec 构建在线自监督自蒸馏的新型学习范式, 利用行为数据有限的用户所产生的自监督信号来提高推荐效果. 相较于 S⁴Rec 模型, 本文模型在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 14.73%, 9.86%.

⑧ SparseEnNet 可生成鲁棒的增强项目来缓解推荐中的数据稀疏性问题, 并搭建一个自训练增强学习模块, 用于捕获相似序列间的一致性关系. 相较于 SparseEnNet 模型, 本文模型在 $HR@N$ 和 $NDCG@N$ 指标上分别最少提升 4.93%, 3.39%.

这些方法虽然在一定程度上提升了推荐效率, 但无法捕捉到多对多的会话间隐藏高阶关系, 且均采用单一的对比学习方式, 模型的泛化能力不足, 很难适应不同领域的推荐任务. 本文模型 TIDA-DSSR 构建了对偶视图编码器, 在不同视角下进行两种对比学习, 在充分挖掘会话间多对多的高阶关系的同时, 提升了模型的泛化能力, 使模型在各个数据集上相较于对比模型均有最优的表现.

综上所述, 相比于其他模型, 本文模型 TIDA-DSSR 在 Beauty, Sports, Home 和 Yelp 上的性能均有较大幅度的提升, 整体平均提升达 5.93%. 本文模型在不同数据集上的表现差异, 既体现了时间间隔数据增强策略的有效性, 也揭示了其在泛化能力方面存在一定的优化空间. 在 Beauty 数据集上, 模型表现显著优于其他模型, 平均提升 11.14%; 在 Sports 和 Home 数据集上分别平均提升 5.03% 和 5.26%. 这些领域的特点是商品品牌集中度较高、属性相似性较强, 用户会话中往往蕴含更多隐藏的高阶关系. 模型通过时间间隔数据增强与对偶视图编码器的结合, 能更好地捕捉这些特性, 充分挖掘会话间的隐藏关系. 然而, Yelp 数据集中的用户行为受到更多外部因素 (如地理位置、社交推荐) 的影响, 时间间隔对用户偏好的表达能力相对较弱, 但本文模型在此数据集上仍能平均提升 2.28%. 为应对不同类型数据对时间间隔敏感性的差异, 本文提出的模型通过自适应调整机制优化时间间隔增强策略, 以适应不同数据集特性. 具体而言, 模型可在训练过程中动态评估时间间隔对推荐性能的贡献, 通过调节时间间隔数据增强模块的权重, 自动适配到不同领域的数据分布. 例如, 在商品属性高度相似的场景 (如 Beauty 数据集), 模型会增强对时间间隔的利用, 以充分挖掘会话内及会话间的高阶关系; 而在属性分布较分散或受外部因素影响较大的场景 (如 Yelp 数据集), 模型则会降低时间间隔的权重, 同时更加注重融合更多上下文特征. 通过这种自适应的时间间隔权重调整, 所提模型 TIDA-DSSR 不仅可保持在多数数据集上的高性能, 还能在多样化场景下展现出更强的泛化能力. 另外, 从指标值的角度看: HR 值与 $NDCG$ 值分别平均提升 5.96% 和 5.89%. 这是因为本文模型构建了更合理的多视角自监督任务框架来提高模型的泛化能力, 并改善因数据增强早期可能带来的无关项目的干扰.

4.4.2 消融实验 (RQ2)

为验证时间间隔数据增强、对偶视图编码、原始会话对比学习这 3 个构件对推荐性能的影响, 构建了 5 种变体模型, 将本文模型 TIDA-DSSR 与它们进行对比, 以证明各构件存在的必要性. 其中, ① Ours-1 不考虑时间信息对会话的数据增强, 而是用原始会话进行编码. 将该变体模型作为对比, 主要考察本文模型的时间间隔数据增强模块的作用; ② Ours-2 仅用 Transformer 编码器进行编码, 不考虑对偶视图编码对会话间信息的挖掘影响. 将该变体模型进行对比, 主要阐明本文模型的对偶视图编码模块的作用; ③ Ours-3 去除了以原始会话编码作为对比学习任务信号的构件, 直接使用数据增强后的会话进行对比学习. 将该变体模型作为对比, 主要验证本文模型的原始会话对比学习模块的合理性; ④ Ours-4 仅保留数据增强模块, 不使用对偶视图编码器进行编码, 也不加入原始会话对比学习. 将该变体模型作为对比, 主要说明本文模型的数据增强模块与其他构件组合的效果如何; ⑤ Ours-5 仅保留对偶视图编码模块, 不实施数据增强, 也不加入原始会话对比学习. 将该变体模型作为对比, 主要阐述对偶视图编码模块与其他构件组合的有效性. 而对于仅保留原始会话对比学习构件这种变体模型无需讨论, 因为仅依靠该构件无法完成自监督任务. 各模型的构件描述情况如表 5 所示.

表 5 各模型构件描述情况

变体模型	时间间隔数据增强	对偶视图编码	原始会话对比学习
Ours-1	×	○	○
Ours-2	○	×	○
Ours-3	○	○	×
Ours-4	○	×	×
Ours-5	×	○	×
TIDA-DSSR	○	○	○

在 Beauty, Sports, Home 及 Yelp 这 4 个公开数据集上进行消融实验, 考虑到变体模型的目的是验证模型构件对性能的影响, 选择 $HR@20$ 和 $NDCG@20$ (下简称 HR 和 $NDCG$) 作为评价指标. 实验结果如表 6 所示.

表 6 消融实验结果对比

模型	Beauty		Sports		Home		Yelp	
	HR	$NDCG$	HR	$NDCG$	HR	$NDCG$	HR	$NDCG$
Ours-1	<u>0.1032</u>	0.0494	0.0589	0.0257	<u>0.0363</u>	0.0179	0.0915	0.0433
Ours-2	0.1024	<u>0.0504</u>	<u>0.0742</u>	<u>0.0286</u>	0.0351	0.0176	<u>0.0983</u>	<u>0.0441</u>
Ours-3	0.0957	0.0497	0.0615	0.027	0.0357	<u>0.0182</u>	0.0867	0.0421
Ours-4	0.0837	0.0483	0.0562	0.0243	0.0343	0.0168	0.0859	0.0427
Ours-5	0.0792	0.0472	0.0658	0.0265	0.0339	0.0163	0.0842	0.0416
TIDA-DSSR	0.1249	0.0568	0.0822	0.0305	0.0412	0.0196	0.1051	0.0472
性能提升 (%)	17.3	11.3	10.9	11.0	13.4	13.5	6.9	7.0

注: 本文模型TIDA-DSSR的实验数据以加粗字体标示, 为便于比较, 利用下划线来突显变体模型中表现最佳的数据, 最后一行给出模型相对于某一最佳变体模型的性能提升情况 (粗体标示)

从表 6 中可看出, 各构件的实验效果分析如下: 时间间隔数据增强构件对模型性能影响较大, 使用原始会话的变体模型 Ours-1 在 HR , $NDCG$ 这 2 个指标上的表现均不如使用时间间隔数据增强的方法; 对偶视图编码构件对于模型性能的影响非常明显, 与仅使用 Transformer 编码器的变体模型 Ours-2 相比, 使用对偶视图的效果有明显提升; 从原始会话对比学习构件的实验结果可看出, 与仅使用数据增强后会话的变体模型 Ours-3 相比, 此构件对模型效果影响最大.

1) 与仅采用原始会话相比, 基于时间信息数据增强后的会话对模型性能影响较大 (变体模型 Ours-1). 为针对现有多数模型未充分利用会话侧信息的问题, 本文模型 TIDA-DSSR 考虑使用基于时间间隔的数据增强模块. 尽管多数用户的兴趣是影响其交互项目的主要因素, 但用户兴趣本身也受时间因素的影响, 所以用户的交互项目也会受时间因素的影响, 且会话内项目间的间隔时长可体现用户兴趣的变化. 由表 6 可看出, 与变体模型 Ours-1 相比, 所提模型性能提升较大, HR 与 $NDCG$ 指标分别平均提升 22.2% 和 13.0%. 其中, 在 HR 指标上至少提升 13.5%, 最多提升可达 39.6%. 由此可知, 本文基于时间间隔信息对会话序列实施数据增强, 丰富会话内信息, 进而细化用户兴趣, 可更有效地提升推荐列表的准确度.

2) 与仅使用 Transformer 编码器相比, 对偶视图编码器对于模型性能的影响明显 (变体模型 Ours-2). 因会话推荐中样本数量少, 且单个会话本身可利用的信息较少, 故如何挖掘会话间的其他信息成为越来越重要的问题. 本文模型 TIDA-DSSR 构建对偶视图编码器, 将会话序列设计成超图, 再使用超图卷积网络编码器将不同会话间联系起来, 实现不同会话间的信息互传递, 充分挖掘不同会话间的相关性, 同时结合 Transformer 编码器, 分别从会话间和会话内提取可利用信息. 由表 6 可看出, 与变体模型 Ours-2 相比, 所提模型性能提升明显, HR 与 $NDCG$ 指标分别平均提升 14.3% 和 9.4%. 其中, 在 HR 指标上至少提升 6.9%, 最多提升可达 22.0%. 由此可知, 对偶视图编码器在多角度挖掘会话间与会话内信息, 提升会话信息的多样性, 在推荐结果上起很大作用.

3) 与仅使用数据增强后会话嵌入的对比学习相比, 原始会话对比学习对于模型效果影响最大 (变体模型 Ours-3). 由于以数据增强后的会话嵌入作为推荐任务的主体, 在模型训练前期对于项目的编码以及项目间相似度计算的准确度不高, 从而易出现会话嵌入受到无关项目影响的现象, 所以本文模型 TIDA-DSSR 同时采用原始会话作为对比学习的样本, 从多角度提取会话信息, 缓解无关项目对于会话嵌入的影响. 这也说明了本文模型各构件之间相互匹配的合理性. 由表 6 可看出, 与变体模型 Ours-3 相比, 所提模型性能提升明显, HR 与 $NDCG$ 指标分别平均提升 25.2% 和 11.8%. 其中, 在 HR 指标上, 至少提升 15.4%, 最多提升可达 33.7%. 由此可知, 加入原始会话对比学习组成的自监督任务框架对提升推荐性能非常有效.

4) 从仅保留时间间隔数据增强构件的变体模型来看, 数据增强构件与其他构件的组合是合理的 (变体模型 Ours-4). 本文模型 TIDA-DSSR 使用基于时间间隔的数据增强模块, 得到增强后的会话信息需进一步处理. 而对偶视图编码模块和自监督任务学习模块可提供处理及融合增强后会话信息的方法. 由表 6 可看出, 所提模型 TIDA-

DSSR 性能提升明显, *HR* 与 *NDCG* 指标分别平均提升 34.5% 和 17.6%. 其中, 在 *HR* 指标上, 至少提升 20.1%, 最多提升可达 49.2%. 与 Ours-4 相比, Ours-2 和 Ours-3 变体模型效果均得到提升. 由此可知, 数据增强模块与其他构件组合对于本文模型性能具有促进作用.

5) 从仅保留对偶视图编码器构件的变体模型来看, 对偶视图编码构件与其他构件的组合是有效的 (变体模型 Ours-5). 本文模型 TIDA-DSSR 使用对偶视图编码器分别从会话内和会话间建模, 这需提供不同信息的会话序列和融合不同角度的信息, 而基于时间间隔的数据增强模块与自监督任务学习模块可较好地完成上述任务. 由表 6 可看出, 与 Ours-5 变体模型相比, 所提模型 TIDA-DSSR 性能提升明显, *HR* 与 *NDCG* 指标分别平均提升 32.2% 和 17.3%. 其中, 在 *HR* 指标上, 至少提升 20.3%, 最多提升可达 57.7%. 与 Ours-5 相比, Ours-1 和 Ours-3 变体模型效果均得到改善. 由此可知, 对偶视图编码器与其他构件的组合对于本文模型性能具有积极作用.

综上分析可知: ① 与仅采用原始会话相比, 经过时间间隔数据增强后的会话充分利用了会话的侧信息, 使得会话嵌入的信息更加丰富, 可更准确地捕捉用户兴趣. ② 利用超图卷积网络编码器提取不同会话间的信息, 再与 Transformer 编码器形成对偶视图对会话序列编码, 使会话嵌入所含信息更加全面. ③ 加入原始会话信息组成多角度的自监督任务, 可有效地缓解模型训练早期数据增强带来的无关项目干扰问题, 提升模型的泛化能力. ④ 时间间隔数据增强构件、对偶视图编码构件与其他构件的组合是合理的, 对于本文模型性能的提升具有非常重要的作用.

4.4.3 构件组合合理性实验 (RQ3)

为验证本文模型中的其中 2 个核心构件 (时间间隔数据增强、对偶视图编码) 的合理性及其独立应用的效果, 本节设计了一组实验, 将这 2 个构件分别加入或替换至 DHCN 与 CoSeRec 这两个经典的对比模型中, 分析各构件给对比模型性能带来的影响, 以评估它们在其他系统中的适用性和通用性. 注意, 因对比模型中信息通道不同, 它们的对比学习框架存在较大差异, 与所提模型的对比学习框架不可兼容, 故在此不讨论替换对比学习框架的情况. 在数据集选择方面, 由于 Sports 数据集在本文所使用的 4 个数据集中特征分布最为均衡, 具有较高的普适性, 因此实验基于 Sports 数据集展开. 在评价指标上, 选取 *HR@20* 和 *NDCG@20* (分别简记为 *HR* 和 *NDCG*) 作为主要衡量标准, 用以综合评估推荐结果的准确性与排序质量. 实验结果展示了各构件单独应用于对比模型后对模型性能的改变情况, 如表 7 所示. 通过对比分析, 判断这些构件是否能够在不同模型中独立发挥作用, 并进一步验证其在提升推荐任务性能方面的有效性.

表 7 构件组合合理性实验结果对比

对比项	时间间隔数据增强				对偶视图编码			
	DHCN		CoSeRec		DHCN		CoSeRec	
	<i>HR</i>	<i>NDCG</i>	<i>HR</i>	<i>NDCG</i>	<i>HR</i>	<i>NDCG</i>	<i>HR</i>	<i>NDCG</i>
原对比模型指标	0.0532	0.0264	0.0629	0.0292	0.0532	0.0264	0.0629	0.0292
加入/替换后指标	0.0517	0.0256	0.0652	0.0316	0.0503	0.0247	0.0602	0.0271
与原对比模型比较 (%)	2.82↓	3.03↓	3.66↑	8.22↑	5.45↓	6.44↓	4.29↓	7.19↓
本模型指标	0.0822	0.0373	0.0822	0.0373	0.0822	0.0373	0.0822	0.0373
与本模型比较 (%)	58.99↓	45.70↓	26.07↓	18.04↓	63.42↓	51.01↓	36.54↓	37.64↓

在引入时间间隔数据增强构件后, 不同模型的表现差异显著. 对于 DHCN 模型, *HR* 指标下降了 2.82%, *NDCG* 指标下降了 3.03%. 这主要是由于 DHCN 模型中的超图节点卷积网络和超边卷积网络在时间间隔数据增强的基础上链接了更多的邻居节点. 然而, 过多的邻居节点信息会稀释目标会话嵌入对自身信息的表达能力, 导致难以精准捕捉用户意图, 从而引发性能下降. 相比之下, 对于 CoSeRec 模型, *HR* 和 *NDCG* 指标分别提升了 3.66% 和 8.22%, 展现了一定的性能改进. 这表明 CoSeRec 模型的框架设计更能有效地利用时间间隔数据增强所带来的信息, 提升推荐质量. 即便如此, 与本文提出的 TIDA-DSSR 模型相比, 其性能仍存在较大的差距, 在 *HR* 和 *NDCG* 指标上分别相差 26.07% 和 18.04%.

在替换对偶视图编码构件后, 不同模型的性能同样受到显著影响. 对于 DHCN 模型, *HR* 和 *NDCG* 指标分别

下降了 5.45% 和 6.44%。这是因为对偶视图编码构件的替换导致 DHCN 模型无法充分利用超边卷积网络所提取的会话间信息,从而使会话嵌入缺乏多样性,进而影响推荐效果。在 CoSeRec 模型中,替换对偶视图编码构件后,HR 和 NDCG 指标分别下降了 4.29% 和 7.19%。其原因在于,CoSeRec 模型的数据增强策略未根据对偶视图编码构件的特点进行有针对性的调整。例如,CoSeRec 中的 Reorder 操作会改变用户会话中项目的顺序,这种改动可能在卷积过程中引入过多无关的邻居节点,干扰会话间信息的提取,从而对推荐结果造成负面影响。

综上所述,本文提出的 TIDA-DSSR 模型突出了时间间隔数据增强、对偶视图编码以及多视角对比学习框架在推荐中互相协同。时间间隔数据增强通过捕捉用户行为的时间动态特征,为模型提供了更丰富的信息表达;对偶视图编码通过整合会话间的信息,显著提升了嵌入的多样性和精度;而多视角对比学习框架则进一步优化了模型的泛化能力。这 3 个核心构件的相互结合,实现了推荐性能提升的最大化。然而,为确保这些构件在不同模型和场景中的兼容性和有效性,所提模型进行了精细的设计和调整,从而使其能在复杂推荐任务中充分发挥作用。

4.4.4 参数敏感度实验 (RQ4)

本文模型 TIDA-DSSR 包含 3 大模块,在时间间隔数据增强模块中,均匀会话占比(记为 σ)是该模块的关键参数;在对偶视图编码模块中,超图卷积网络层数(记为 θ)是对该模块影响最大的参数;在自监督学习任务模块中,超图卷积权重(记为 λ_1)和原始会话对比学习权重(记为 λ_2)是最关键的 2 个参数。故本节通过实验分析这 4 个参数对模型性能的影响。选取 HR@10 和 NDCG@10(简记为 HR 和 NDCG)作为评价指标。

1) 均匀会话占比 σ 的影响

σ 是决定均匀会话在全体会话中占比的参数,体现会话数据增强的范围大小。 σ 过大意味着会话中均匀会话占比较大,则需数据增强的会话占比就少,从而缩小数据增强的作用范围; σ 过小意味着需数据增强的会话占比较多,使得较为完整的会话也会进行数据增强,很难体现用户的真实偏好。 σ 的值从集合{0, 0.1, 0.2, 0.3, 0.4, 0.5}中选取,各指标值随 σ 的变化情况如图 2 所示。

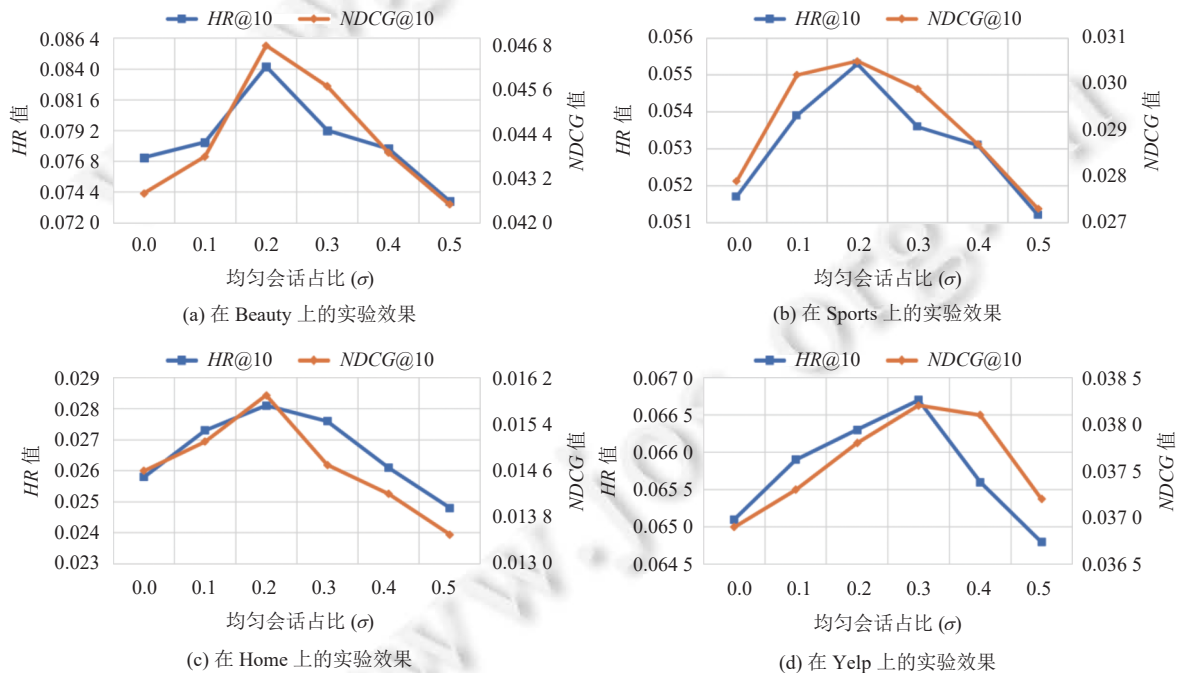


图 2 参数 σ 的影响

从图 2 中可看出,随着 σ 的增大,模型的性能先上升后下降。当 σ 值为 0.2 时,模型在 Beauty、Sports 以及 Home 数据集上的性能达到最优。而在 Yelp 数据集上,当 σ 值为 0.3 时,模型性能为最优。推荐效果先增后减的这

种变化趋势表明, 当 σ 取值为 0.2 时, 模型对于用户会话的数据增强范围最合适.

可见, 在此参数实验中, 大多数数据集下 σ 的最优取值相同, 我们认为这是合理的. 因为大多数数据集中用户交互的数据对于捕捉用户兴趣偏好是不足的, 需对大部分会话实施数据增强以更准确地捕捉用户偏好.

2) 超图卷积网络层数 θ 的影响

θ 是决定超图卷积网络内相关会话传播深度的因素, 体现对隐藏高阶关系的敏感程度. θ 过大会导致当前会话嵌入中融合相距过远的无关会话信息, 会对隐藏高阶关系捕捉不准确, 从而影响推荐效果. θ 过小会使获取的会话嵌入中融合的相似会话信息过少, 对隐藏高阶关系挖掘不足, 推荐效果反而接近卷积前的效果. θ 的值从集合 $\{0, 1, 2, 3, 4\}$ 中选取, 各指标值随 θ 的变化情况如图 3 所示.

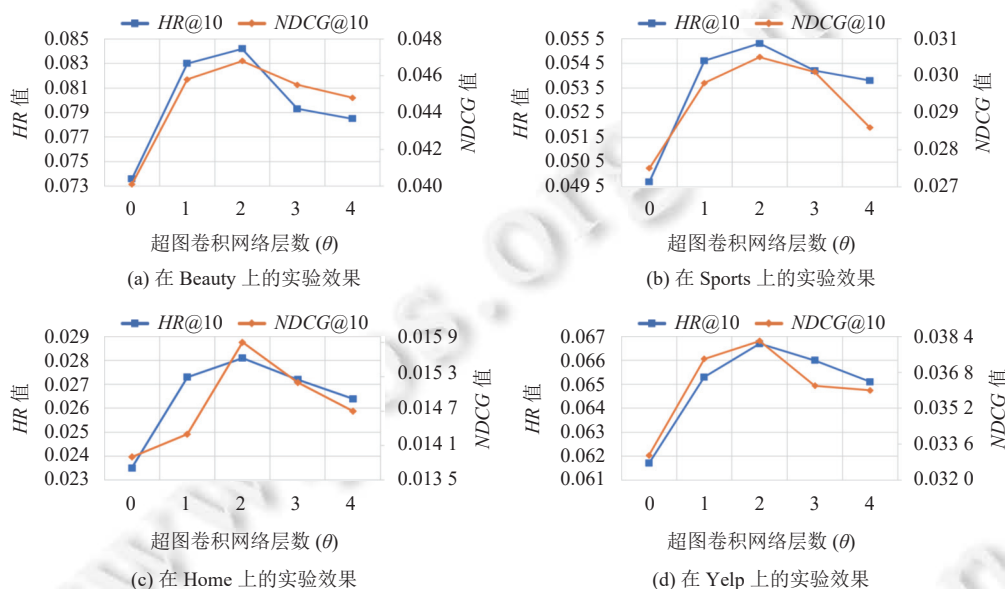


图 3 参数 θ 的影响

从图 3 中可看出, 当 θ 由 0 变为 1 时, 本文模型性能提升明显, 说明对偶视图编码器中的超图卷积网络能有效地捕捉会话间的隐藏高阶关系. 第 4.4.2 节消融实验中的变体模型 Ours-2 可看作 θ 为 0 的情况, 此时各指标值最小, 在此参数敏感度实验中推荐效果最差. 随着 θ 的增大, 模型的性能先上升后下降. 当 θ 值为 2 时, 模型在各数据集上的性能达到最优. 推荐效果先增后减的这种变化趋势, 表明当 θ 取值为 2 时, 模型对于用户会话间隐藏高阶关系捕捉能力最强.

可见, 在此参数实验中, 各数据集下 θ 的最优取值相同, 我们认为这是合理的. 这是因为当用户的相似会话相距过远时, 有效的会话间信息会减少. 换言之, 模型捕捉距离 2 步以内的相似会话关系是最合理的, 这也与我们对实际的认知相吻合.

3) 超图卷积权重 λ_1 的影响

λ_1 是控制会话嵌入中会话间信息融合多少的因素, 体现相关会话对于当前会话的影响程度. λ_1 过大会导致会话嵌入中掺杂太多相关会话的信息, 使得当前会话的信息不够明确, 从而无法捕捉用户真正的兴趣偏好; 而 λ_1 过小则会导致会话嵌入丧失相关会话的信息, 丢失重要的会话间信息. 因此, λ_1 在很大程度上影响对会话间高阶隐藏关系的捕捉. λ_1 的值从集合 $\{1, 0.1, 0.001, 0.0001\}$ 中选取, 各指标值随 λ_1 的变化情况如图 4 所示.

从图 4 中可看出, 随着 λ_1 增大, 所提模型 TIDA-DSSR 在不同数据集下的表现有较明显差异: 在 Beauty 数据集下, λ_1 值为 0.1 时, HR 和 NDCG 两个指标值均达到最大. 在 Sports 数据集下, λ_1 值为 0.001 时, 各指标值均最大. 在 Home 数据集下, λ_1 值为 0.1 时, HR 指标值达到最大, 而 NDCG 指标在 λ_1 取 0.001 时效果最好. 在 Yelp 数据集

下, 两个指标 HR 和 $NDCG$ 分别在 λ_1 取 0.001、0.0001 时获得最大值.

可见, 在不同数据集上, 参数 λ_1 的最优取值有所不同, 我们认为这是合理的. 因为每个数据集的特点不同, 会话间信息密度及重要性也不同, 展现的隐藏高阶关系频度也不一样, 只有当 λ_1 的取值与数据集的属性相匹配时, 才能使模型达到最佳效果.

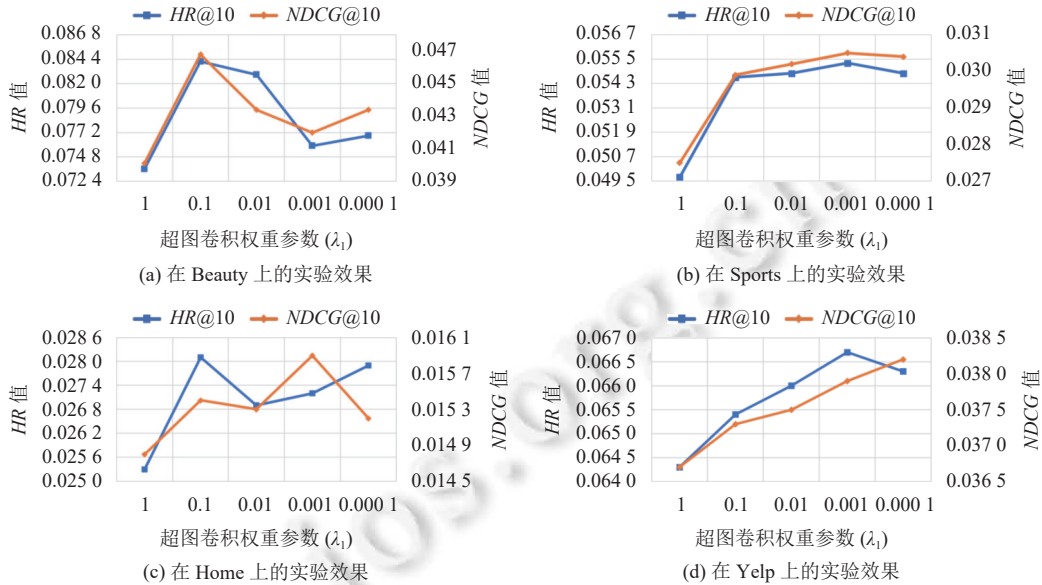


图4 参数 λ_1 的影响

4) 原始会话对比学习权重 λ_2 的影响

λ_2 是控制原始会话对比学习占联合学习比重的参数, 体现联合学习受到原始会话影响的大小. λ_2 值过大会导致会话嵌入受原始会话影响太多, 这会降低相关联会话间信息的影响; λ_2 值过小会导致会话嵌入易受到无关项目的影响, 从而降低推荐效果. λ_2 的值从集合 $\{1, 0.1, 0.001, 0.0001\}$ 中选取, 各指标值随 λ_2 的变化情况如图 5 所示.

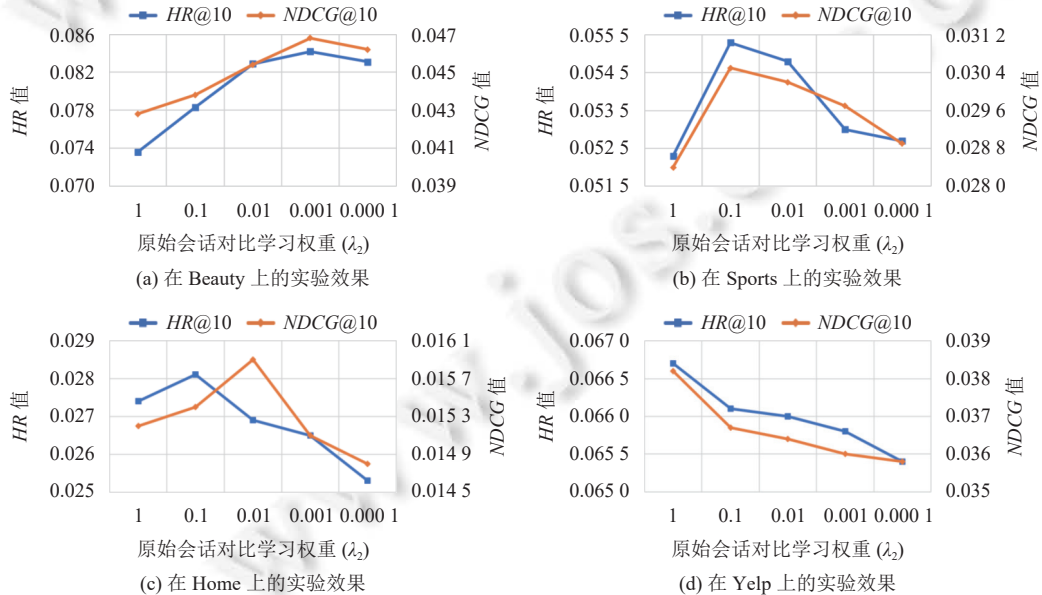


图5 参数 λ_2 的影响

从图5中可看出, 随着 λ_2 增大, 所提模型 TIDA-DSSR 在不同数据集下的表现有明显差异: 在 Beauty 数据集下, λ_2 值为 0.001 时, HR 和 $NDCG$ 两个指标值均最大; 在 Sports 数据集下, λ_2 值为 0.1 时, 各指标值均最大; 在 Home 数据集下, λ_2 值为 0.1 时, HR 指标值达到最大, 而 $NDCG$ 指标在 λ_2 取 0.01 时效果最好; 在 Yelp 数据集下, 各指标在 λ_2 取值为 1 时达到最大。

可见, 在不同数据集上, 参数 λ_2 的最优取值有所不同, 我们认为这是合理的。因为当数据集中的项目种类越多, 用户交互序列中出现无关项目的可能性就越大, 此时需加大原始会话对比学习权重 λ_2 的值, 使得会话嵌入受到无关项目的影响降低。即, 只有当 λ_2 的取值与数据集相匹配时, 才能使模型效果最佳。

4.4.5 模型复杂度分析 (RQ5)

所提模型 TIDA-DSSR 利用 HGCN 以及 Transformer 等技术实现推荐任务, 与现有模型相比, 推荐步骤相对增多, 尽管复杂度随之增加, 但模型的推荐效果得到较大提升。下面是相关模型的时间复杂度分析。

1) 本文模型 TIDA-DSSR 的时间复杂度, 主要表现在超图卷积网络、Transformer 编码器模块。超图卷积网络部分的时间复杂度为 $O(NEFL)$, 其中 N 为超图的节点, E 为超图中的超边数, F 为特征维度, L 为超图卷积层数; Transformer 编码器由多个相同的模块串联而成, 主要包括自注意力层和前馈全连接层, 它的时间复杂度主要取决于自注意力层, 时间复杂度为 $O(Sn^2F)$, 其中 S 为模块数, n 为输入的序列长度。因此, 本文模型的时间复杂度大致为 $O(NEFL + Sn^2F)$ 。

2) 就对比模型 LightGCN (该模型较经典, 在非自监督对比模型中复杂度最小) 来说, 它的时间主要消耗在图卷积层的动态扩散过程中。因此, 时间复杂度大致为 $O(MDB)$, 其中 M 是交互图节点数, D 为扩散深度, B 为边数。

3) 就对比模型 DHCN (该模型较新, 与本文结构最为相似) 来说, 它的时间主要消耗在超图节点卷积网络和超边卷积网络中。因此, 时间复杂度大致为 $O(NEFL + NEFP)$, 其中 P 为超边卷积网络层数。

4) 就对比模型 TiCoSeRec (该模型较新, 在自监督对比模型中推荐性能最好) 来说, 它的时间主要消耗在相关交互矩阵计算。因此, 时间复杂度大致为 $O(HnF + M^3)$, 其中 H 为隐藏层大小。

5) 就对比模型 SSDRec (该模型最新, 在自监督对比模型中复杂度最小) 来说, 它的时间主要消耗在全局关系编码和序列自增强过程。因此, 时间复杂度大致为 $O(UVn + Vn)$, 其中 U 为用户数, V 为项目数。

综上分析, 各模型的时间复杂度大小关系为: $DHCN > TIDA-DSSR > TiCoSeRec > LightGCN > SSDRec$ 。从时间复杂度上来看, 我们的模型 TIDA-DSSR 较对比模型 DHCN 更低, 并不是最高的, 尽管要略高于其他部分对比模型, 但结合推荐效果来说, 其设计思想依然是合理的、有意义的。

1) 与模型 DHCN 相比, 所提模型 TIDA-DSSR 的时间复杂度相对较低, 训练所需的资源较少, 且利用用户交互项目序列中的邻居用户信息对可能的隐藏高阶关系进一步挖掘, 使模型可更准确地预测结果, 在 HR 和 $NDCG$ 指标上分别最少提升 16.61%、20.41% (见表4), 在增强推荐效果方面更出色。

2) 与模型 LightGCN 相比, 尽管所提模型 TIDA-DSSR 时间复杂度相对较高, 但推荐效果比 LightGCN 在 HR 和 $NDCG$ 指标上分别最少提升 23.06%、16.46% (见表4), 推荐性能提升明显。

3) 与模型 TiCoSeRec 相比, 尽管所提模型 TIDA-DSSR 的时间复杂度相对较高, 但 TiCoSeRec 未充分利用会话间信息, 而本文模型从多角度考虑了会话间与会话内的信息, 丰富会话嵌入的多样性, 提高了推荐准确率, 在 HR 和 $NDCG$ 指标上分别最少提升 1.37%、2.83% (见表4), 推荐性能更占优势。

4) 与模型 SSDRec 相比, 尽管所提模型 TIDA-DSSR 的时间复杂度相对较高, 但 SSDRec 在项目嵌入建模时考虑的信息过于单一, 而本文模型充分挖掘时间信息, 并从不同视角建模项目嵌入, 使推荐更精准, 比 SSDRec 在 HR 和 $NDCG$ 指标上分别最少提升 53.55%、34.75% (见表4), 推荐性能优势显著。

根据分析, 我们的模型 TIDA-DSSR 在计算复杂度方面是可接受的, 并通过合理的架构设计优化了计算资源的利用效率。具体而言, 模型的核心组件在参数规模和计算时间上均得到了良好的平衡, 可适配常见的深度学习硬件环境。目前市面上主流用于深度学习的设备 (如高性能 GPU 或 TPU 集群) 完全能够满足 TIDA-DSSR 模型的计算需求, 该模型在实际部署时是可行的。此外, 所提模型在处理推荐任务时展现出较强的性能优势, 能够有效地缓解数据稀疏性、挖掘高阶关系以及提升推荐准确性等。

4.4.6 模型框架分析 (RQ6)

本文模型 TIDA-DSSR 通过对会话内和会话间的特征建模, 已初步实现了对不同尺度信息的捕捉. 为进一步提升对复杂推荐场景的适应性, 可将对比学习框架扩展为多尺度结构. 在局部层面上, 模型可深度挖掘单一会话序列中行为间的时序依赖性; 在全局层面上, 模型可通过构建跨会话的超图结构, 提取用户间的协同行为或项目间的语义相似性. 通过结合局部与全局特征的对比, 模型能够更好地平衡短期兴趣与长期兴趣的关联, 从而进一步提升其泛化能力. 此外, 多视角 (如用户与项目、时间与行为) 的对比学习丰富了模型在多样化场景下的表现能力, 有助于更精准地捕捉复杂推荐场景中的隐式关系.

所提模型 TIDA-DSSR 当前利用会话内时序信息和会话间高阶关系进行对比学习, 但对高阶关系的建模仍可进一步拓展. 通过显式构建数据中的高阶关系图 (如跨会话项目间的隐式语义关系或用户间的协同行为), 可设计关系更复杂的对比目标. 具体而言, 在关系图中, 节点可代表用户、会话或项目, 边则表示它们之间的隐式关联 (如语义相似性或协作行为). 通过在对比学习框架中融入这些高阶关系, TIDA-DSSR 模型能够捕捉更加丰富的上下文信息, 挖掘不同数据分布中的潜在信号, 提升跨领域场景下的推荐性能.

此外, 通过引入知识图谱, 将项目之间的显式关联 (如类别、品牌) 与隐式关联 (如共现关系) 融合到对比学习中, 可帮助模型更好地理解复杂数据分布. 而且, 将上下文信息 (如地理位置、时间、设备类型) 融入到对比任务中, 能够构建场景感知能力更强的对比目标.

故本文提出的对比学习框架灵活性较强, 可有效地应对复杂的推荐场景. 通过多角度的视图构建和改进后的数据增强策略, 该框架能够适配多种类的数据分布和应用场景, 展现出较强的泛化能力与适应性.

4.4.7 案例分析

上述实验结果表明, 所提模型 TIDA-DSSR 可有效地提高推荐性能. 下面我们进一步利用 Sports 数据集中的数据进行案例分析, 该数据集的项目种类集中度相对较高, 其用户会话更易呈现高阶隐藏关系. 另外, 作为深度学习会话推荐中的较经典模型, DHCN 能够动态捕捉会话间高阶关系, 并将多对多的高阶关系进行具体化, 故使用模型 DHCN 与所提模型 TIDA-DSSR 来预测会话中用户的下一行为并加以对比, 以直观地展示 TIDA-DSSR 的有效性. 具体示例如图 6 所示. 为便于阐述, 这里仅展示部分会话序列.



图 6 案例分析示例

在会话 1 与会话 2 的情况下, 我们可看到, TIDA-DSSR 模型能正确地预测会话 3 中用户下一行为是购买“护膝”. 这是因为 TIDA-DSSR 模型中数据增强模块与对偶视图编码器模块相结合能提取会话间多种相似信息 (即, 会话 2 中的“篮球”与会话 3 中的“足球”相似, 会话 2 中的“篮球服”与会话 3 中的“足球服”相似), 进而捕捉会话间的高阶隐藏关系, 判断出会话 2 与会话 3 相似度更高, 从而根据会话 2 进行预测. 而 DHCN 模型通过会话间的相同项目 (即, 会话 1 与会话 3 中有相同项目“运动水壶”) 错误地将会话 1 与会话 3 联系起来, 会受到整体会话不相似方面的影响, 从而预测不合适的项目“运动头巾”. 可见, 本文模型 TIDA-DSSR 在捕捉会话间的隐藏高阶关系上更有优势.

5 总结与下一步工作

本文提出一种结合时间间隔数据增强的对偶视图自监督会话推荐模型 TIDA-DSSR, 利用时间间隔数据增强充实会话内信息, 更合理地建模用户兴趣偏好, 提高推荐准确度; 设计对偶视图编码器, 利用超图卷积网络和 Transformer 编码器从不同角度对用户会话进行建模, 结合时间间隔数据增强模块可更全面、细致地捕捉用户会话间的隐藏高阶关系, 丰富会话信息多样性; 构建一种新的自监督学习框架, 在对比学习中加入原始会话信息作为辅助任务, 使得数据增强后的会话信息更全面的同时, 还可减少数据增强早期可能带来无关项目的影响, 提升模型泛化能力。

所提模型 TIDA-DSSR 无论是与经典模型还是最新模型对比, 在 2 个常用推荐指标上均有明显提升, 但依然存在一些待完善之处。在接下来工作中将继续探讨会话中不同关系建模的可能性, 分析这些关系的内在影响规律, 据此挖掘更多影响模型性能的因素, 从而进一步优化模型。

References:

- [1] Wang SJ, Pasi G, Hu L, Cao LB. The era of intelligent recommendation: Editorial on intelligent recommendation with advanced AI and learning. *IEEE Intelligent Systems*, 2020, 35(5): 3–6. [doi: [10.1109/MIS.2020.3026430](https://doi.org/10.1109/MIS.2020.3026430)]
- [2] 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. 1–10.
- [3] Qiu RH, Huang Z, Li JJ, Yin HZ. Exploiting cross-session information for session-based recommendation with graph neural networks. *ACM Trans. on Information Systems (TOIS)*, 2020, 38(3): 22. [doi: [10.1145/3382764](https://doi.org/10.1145/3382764)]
- [4] Bu JJ, Tan SL, Chen C, Wang C, Wu H, Zhang LJ, He XF. Music recommendation by unified hypergraph: Combining social media information and music content. In: *Proc. of the 18th ACM Int'l Conf. on Multimedia*. Firenze: ACM, 2010. 391–400. [doi: [10.1145/1873951.1874005](https://doi.org/10.1145/1873951.1874005)]
- [5] Yadati N, Nimishakavi M, Yadav P, Nitin V, Louis A, Talukdar P. HyperGCN: A new method of training graph convolutional networks on hypergraphs. In: *Proc. of the 33rd Int'l Conf. on Neural Information Processing Systems*. Curran Associates Inc., 2019. 135.
- [6] Jiang JW, Wei YX, Feng YF, Cao JX, Gao Yue. Dynamic hypergraph neural networks. In: *Proc. of the 28th Int'l Joint Conf. on Artificial Intelligence*. Macao: ijcai.org, 2019. 2635–2641. [doi: [10.24963/ijcai.2019/366](https://doi.org/10.24963/ijcai.2019/366)]
- [7] Bandyopadhyay S, Das K, Narasimha Murty M. Line hypergraph convolution network: Applying graph convolution for hypergraphs. *arXiv:2002.03392*, 2020.
- [8] Xia X, Yin HZ, Yu JL, Wang QY, Cui LZ, Zhang XL. Self-supervised hypergraph convolutional networks for session-based recommendation. In: *Proc. of the 35th AAAI Conf. on Artificial Intelligence*. AAAI, 2021. 4503–4511. [doi: [10.1609/aaai.v35i5.16578](https://doi.org/10.1609/aaai.v35i5.16578)]
- [9] 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*. Kuala Lumpur: IEEE, 2022. 1259–1273. [doi: [10.1109/ICDE53745.2022.00099](https://doi.org/10.1109/ICDE53745.2022.00099)]
- [10] Wang W, Zhang W, Liu SK, Liu Q, Zhang B, Lin LY, Zha HY. Beyond clicks: Modeling multi-relational item graph for session-based target behavior prediction. In: *Proc. of the 2020 Web Conf*. Taipei: ACM, 2020. 3056–3062. [doi: [10.1145/3366423.3380077](https://doi.org/10.1145/3366423.3380077)]
- [11] Wang ZY, Wei W, Cong G, Li XL, Mao XL, Qiu MH. Global context enhanced graph neural networks for session-based recommendation. In: *Proc of the 43rd Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval*. ACM, 2020. 169–178. [doi: [10.1145/3397271.3401142](https://doi.org/10.1145/3397271.3401142)]
- [12] Huang ZH, Lin XL, Sun LS, Tang Y, Chen YW. Feature augmentation based graph neural recommendation method in session scenarios. *Chinese Journal of Computers*, 2022, 45(4): 766–780 (in Chinese with English abstract). [doi: [10.11897/SP.J.1016.2022.00766](https://doi.org/10.11897/SP.J.1016.2022.00766)]
- [13] Zhang S, Gao M, Wen JH, Xiong QY, Tang X. Self-supervised learning for alleviating popularity bias in recommender systems. *Acta Electronica Sinica*, 2022, 50(10): 2361–2371 (in Chinese with English abstract). [doi: [10.12263/DZXB.20210443](https://doi.org/10.12263/DZXB.20210443)]
- [14] Yao TS, Yi XY, Cheng DZ, Yu F, Chen T, Menon A, Hong LC, Chi EH, Tjoa S, Kang JQ, Ettinger E. Self-supervised learning for large-scale item recommendations. In: *Proc. of the 30th ACM Int'l Conf. on Information & Knowledge Management*. Queensland: ACM, 2021. 4321–4330. [doi: [10.1145/3459637.3481952](https://doi.org/10.1145/3459637.3481952)]
- [15] Zhou K, Wang H, Zhao WX, Zhu YT, Wang SR, Zhang FZ, Wang ZY, Wen JR. S3-Rec: Self-supervised learning for sequential recommendation with mutual information maximization. In: *Proc. of the 29th ACM Int'l Conf. on Information & Knowledge Management*. ACM, 2020. 1893–1902. [doi: [10.1145/3340531.3411954](https://doi.org/10.1145/3340531.3411954)]
- [16] Qiu RH, Huang Z, Yin HZ, Wang ZJ. Contrastive learning for representation degeneration problem in sequential recommendation. In:

- Proc. of the 15th ACM Int'l Conf. on Web Search and Data Mining. ACM, 2022. 813–823. [doi: [10.1145/3488560.3498433](https://doi.org/10.1145/3488560.3498433)]
- [17] Cai XH, Huang C, Xia LH, Ren XB. LightGCL: Simple yet effective graph contrastive learning for recommendation. arXiv:2302.08191, 2023.
 - [18] Liu ZW, Chen YJ, Li J, Yu PS, McAuley J, Xiong CM. Contrastive self-supervised sequential recommendation with robust augmentation. arXiv:2108.06479, 2021.
 - [19] Dang YZ, Yang EN, Guo GB, Jiang LY, Wang XW, Xu XX, Sun QH, Liu H. Uniform sequence better: Time interval aware data augmentation for sequential recommendation. In: Proc. of the 37th AAAI Conf. on Artificial Intelligence. Washington: AAAI, 2023. 4225–4232. [doi: [10.1609/aaai.v37i4.25540](https://doi.org/10.1609/aaai.v37i4.25540)]
 - [20] Fang Y, Tan Z, Chen ZY, Xiao WD, Zhang LL, Tian F. Contrastive meta-learning on heterogeneous information networks for cold-start recommendation. Ruan Jian Xue Bao/Journal of Software, 2023, 34(10): 4548–4564 (in Chinese with English abstract). <http://www.jos.org.cn/1000-9825/6886.htm> [doi: [10.13328/j.cnki.jos.006886](https://doi.org/10.13328/j.cnki.jos.006886)]
 - [21] Qian ZS, Xiao SL, Zhu H, Wang XW, Liu JP. Long-tail recommendation model utilizing GRU dual-branch information col-laboration enhancement. Journal of Frontiers of Computer Science & Technology, 2025, 19(2): 476–489 (in Chinese with English abstract). [doi: [10.3778/j.issn.1673-9418.2402052](https://doi.org/10.3778/j.issn.1673-9418.2402052)]
 - [22] Xin X, Yang L, Zhao ZQ, Ren PJ, Chen ZM, Ma J, Ren ZC. On the effectiveness of unlearning in session-based recommendation. In: Proc. of the 17th ACM Int'l Conf. on Web Search and Data Mining. Merida: ACM, 2024. 855–863. [doi: [10.1145/3616855.3635823](https://doi.org/10.1145/3616855.3635823)]
 - [23] Zhang C, Han QL, Chen R, Zhao XY, Tang P, Song HT. SSDRec: Self-augmented sequence denoising for sequential recommendation. In: Proc. of the 40th IEEE Int'l Conf. on Data Engineering. Utrecht: IEEE, 2024. 803–815. [doi: [10.1109/ICDE60146.2024.00067](https://doi.org/10.1109/ICDE60146.2024.00067)]
 - [24] Feng YF, You HX, Zhang ZZ, Ji RR, Gao Y. Hypergraph neural networks. In: Proc. of the 33rd AAAI Conf. on Artificial Intelligence. Honolulu: AAAI, 2019. 3558–3565. [doi: [10.1609/aaai.v33i01.33013558](https://doi.org/10.1609/aaai.v33i01.33013558)]
 - [25] Wang JL, Ding KZ, Hong LJ, Liu H, Caverlee J. Next-item recommendation with sequential hypergraphs. In: Proc. of the 43rd Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval. ACM, 2020. 1101–1110. [doi: [10.1145/3397271.3401133](https://doi.org/10.1145/3397271.3401133)]
 - [26] Zhao S, Wei W, Liu YF, Wang ZY, Li WD, Mao XL, Zhu S, Yang MH, Wen ZJ. Towards hierarchical policy learning for conversational recommendation with hypergraph-based reinforcement learning. In: Proc. of the 32nd Int'l Joint Conf. on Artificial Intelligence. Macao: ijcai.org, 2023. 273. [doi: [10.24963/ijcai.2023/273](https://doi.org/10.24963/ijcai.2023/273)]
 - [27] Yu YL, Yang EN, Guo GB, Jiang LY, Wang XW. Handling information loss of graph neural networks for session-based recommendation. In: Proc. of the 32nd Int'l Joint Conf. on Artificial Intelligence. 2023. 2415–2422.
 - [28] Yang LW, Wang SJ, Tao YZ, Sun JK, Liu XL, Yu PS, Wang TQ. DGRec: Graph neural network for recommendation with diversified embedding generation. In: Proc. of the 16th Int'l Conf. on Web Search and Data Mining. Singapore: ACM, 2023. 661–669. [doi: [10.1145/3539597.3570472](https://doi.org/10.1145/3539597.3570472)]
 - [29] Zhang Q, Wu B, Sun ZC, Ye YD. Fine-grained modeling of user interests for sequential recommendation. Scientia Sinica Informationis, 2022, 52(10): 1775–1791 (in Chinese with English abstract). [doi: [10.1360/SSI-2021-0026](https://doi.org/10.1360/SSI-2021-0026)]
 - [30] Yu JL, Yin HZ, Xia X, Chen T, Cui LZ, Nguyen QVH. Are graph augmentations necessary?: Simple graph contrastive learning for recommendation. In: Proc. of the 45th Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval. Madrid: ACM, 2022. 1294–1303. [doi: [10.1145/3477495.3531937](https://doi.org/10.1145/3477495.3531937)]
 - [31] 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.
 - [32] He XN, Deng K, Wang X, Li Y, Zhang YD, Wang M. LightGCN: Simplifying and powering graph convolution network for recommendation. In: Proc. of the 43rd Int'l ACM SIGIR Conf. on Research and Development in Information Retrieval. ACM, 2020. 639–648. [doi: [10.1145/3397271.3401063](https://doi.org/10.1145/3397271.3401063)]
 - [33] Liu Q, Zeng YF, Mokhosi R, Zhang HB. STAMP: Short-term attention/memory priority model for session-based recommendation. In: Proc. of the 24th ACM SIGKDD Int'l Conf. on Knowledge Discovery & Data Mining. London: ACM, 2018. 1831–1839. [doi: [10.1145/3219819.3219950](https://doi.org/10.1145/3219819.3219950)]
 - [34] Sun F, Liu J, Wu J, Pei CH, Lin X, Ou WW, Jiang P. BERT4Rec: Sequential recommendation with bidirectional encoder representations from Transformer. In: Proc. of the 28th ACM Int'l Conf. on Information and Knowledge Management. Beijing: ACM, 2019. 1441–1450. [doi: [10.1145/3357384.3357895](https://doi.org/10.1145/3357384.3357895)]
 - [35] Wei SW, Wu ZW, Li X, Wu QT, Zhang ZQ, Zhou J, Gu LH, Gu JJ. Leave no one behind: Online self-supervised self-distillation for sequential recommendation. In: Proc. of the 2024 ACM Web Conf. Singapore: ACM, 2024. 3767–3776. [doi: [10.1145/3589334.3645590](https://doi.org/10.1145/3589334.3645590)]
 - [36] Chen JY, Zou GX, Zhou P, Wu YR, Chen ZH, Su HC, Wang H, Gong ZG. Sparse enhanced network: An adversarial generation method for robust augmentation in sequential recommendation. In: Proc. of the 38th AAAI Conf. on Artificial Intelligence. Vancouver: AAAI,

2024. 8283–8291. [doi: [10.1609/aaai.v38i8.28669](https://doi.org/10.1609/aaai.v38i8.28669)]

附中文参考文献:

- [12] 黄震华, 林小龙, 孙圣力, 汤庸, 陈运文. 会话场景下基于特征增强的图神经推荐方法. 计算机学报, 2022, 45(4): 766–780. [doi: [10.11897/SP.J.1016.2022.00766](https://doi.org/10.11897/SP.J.1016.2022.00766)]
- [13] 张帅, 高旻, 文俊浩, 熊庆宇, 唐旭. 基于自监督学习的去流行度偏差推荐方法. 电子学报, 2022, 50(10): 2361–2371. [doi: [10.12263/DZXB.20210443](https://doi.org/10.12263/DZXB.20210443)]
- [20] 方阳, 谭真, 陈子阳, 肖卫东, 张玲玲, 田锋. 用于冷启动推荐的异质信息网络对比元学习. 软件学报, 2023, 34(10): 4548–4564. <http://www.jos.org.cn/1000-9825/6886.htm> [doi: [10.13328/j.cnki.jos.006886](https://doi.org/10.13328/j.cnki.jos.006886)]
- [21] 钱忠胜, 肖双龙, 朱辉, 王晓闻, 刘金平. 利用 GRU 双分支信息协同增强的长尾推荐模型. 计算机科学与探索, 2025, 19(2): 476–489. [doi: [10.3778/j.issn.1673-9418.2402052](https://doi.org/10.3778/j.issn.1673-9418.2402052)]
- [29] 张麒, 吴宾, 孙中川, 叶阳东. 细粒度建模用户兴趣的序列化推荐方法. 中国科学: 信息科学, 2022, 52(10): 1775–1791. [doi: [10.1360/SSI-2021-0026](https://doi.org/10.1360/SSI-2021-0026)]



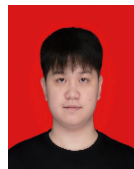
钱忠胜(1977—), 男, 博士, 教授, 博士生导师, CCF 高级会员, 主要研究领域为软件工程, 机器学习, 智能化软件.



范赋宇(2002—), 男, 硕士生, 主要研究领域为软件工程, 机器学习.



万子琰(1999—), 男, 硕士生, 主要研究领域为软件工程, 机器学习.



付庭峰(1999—), 男, 硕士生, 主要研究领域为软件工程, 机器学习.