

Semantic-Aware and Attention Mechanism Based Spam Detection Approach

ZHANG Ling^{1*}, ZHENG Ying¹, SONG Wanli²

(1. Information Management Center, Nanjing Vocational Institute of Transport Technology, Nanjing 211188, China; 2. School of Information Engineering & School of Artificial Intelligence, Nanjing Xiaozhuang University, Nanjing 211171, China)

Abstract: Comment spams pose a significant threat to the reputation of e-commerce platforms. Existing methods based on graph neural networks for detecting spam comments face issues such as data sparsity, complex node relationships, and incomplete characterizations, which affect the detection performance. To address these problems, this paper proposes a graph neural network-based spam detection method, named GNNSD. The proposed method introduces a semantic-aware node enhancement mechanism to generate semantically relevant neighbors for minority class nodes to alleviate the performance decline caused by class imbalance and insufficient features. Further, GNNSD employs multi-head relation-aware attention and semantic-relation dual attention mechanisms to comprehensively capture the interaction patterns and fine-grained dependency features between nodes, and adopts cross-layer residual connection and contrastive learning strategies to enhance feature reuse and increase the discriminative power of node embeddings. Experiments are conducted on benchmark datasets, and the results show that GNNSD outperforms existing mainstream methods in terms of AUC, Micro- F_1 , and Macro- F_1 metrics, achieves better detection performance in the case of class imbalance, and has a good potential in the field of spam detection.

Highlights:

1. A semantic-aware node augmentation mechanism is designed to enhance the feature expression capabilities of minority category nodes, providing a new solution to the class imbalance problem in spam detection.
2. The combination of a multi-head relationship-aware attention mechanism and a semantic-relation attention mechanism enables fine-grained modeling of the complex interaction relationships among comments, users, and products, which comprehensively captures the multi-dimensional heterogeneous information in the e-commerce comment system and significantly improves the detection performance.

Key words: spam detection; heterogeneous graph neural network; data sparsity; representation learning; multi-attention mechanism

Foundation items: China Transportation Education Research Association Key Project (2022—2024)(No. JT2022ZD043); Nanjing Vocational Institute of Transport Technology Research Project(No. JZ2207); Key Special Project of Jiangsu Province 2026: New Ecosystem of Higher Education Empowered by Artificial Intelligence(No. 2026XST041).

Received: 2025-11-25; **Revised:** 2026-01-23

***Corresponding author, E-mail:** 4033830@qq.com.

融合语义感知与注意力机制的垃圾信息检测方法

张 鸽¹, 郑 莹¹, 宋万里²

(1. 南京交通职业技术学院信息化管理中心, 南京 211188; 2. 南京晓庄学院信息工程学院、人工智能学院, 南京 211171)

摘 要: 垃圾信息是影响电商平台信誉的重要威胁之一。面对类别不平衡的评论数据, 现有的基于神经网络的垃圾信息检测方法面临数据稀疏、节点关系复杂且刻画不全面等问题, 影响了检测性能。为了应对上述问题, 本文提出了一种基于图神经网络的垃圾信息检测方法(Graph neural network-based spam detection, GNNSD)。该方法引入语义感知的节点增强机制, 为少数类节点生成语义相关的邻居, 以缓解类别不平衡导致的性能下降。同时, 采用多头关系感知注意力与语义-关系双重注意力机制, 捕捉节点间的交互模式和细粒度的依赖特征, 解决特征信息利用不充分问题。进一步地, 采用跨层残差连接方式和对比学习策略, 强化特征复用与增加节点嵌入的区分性。在基准数据集上开展了实验验证, 结果表明 GNNSD 在 AUC、Micro- F_1 和 Macro- F_1 指标上优于现有主流方法, 在类别不平衡条件下取得了更佳的检测性能。

关键词: 垃圾信息检测; 异构图神经网络; 数据稀疏; 表示学习; 多重注意力机制

中图分类号: TP391

文献标志码: A

引用格式: 张鸽, 郑莹, 宋万里. 融合语义感知与注意力机制的垃圾信息检测方法[J]. 数据采集与处理, 2026, 41(4): 1226-1238. ZHANG Ling, ZHENG Ying, SONG Wanli. Semantic-aware and attention mechanism based spam detection approach[J]. Journal of Data Acquisition and Processing, 2026, 41(4): 1226-1238.

引 言

在互联网时代, 垃圾信息已成为影响用户体验和平台信任度的主要威胁之一。特别是在电商平台, 垃圾评论的泛滥不仅影响用户的购物和决策体验, 而且降低商户声誉, 进而影响平台生态。垃圾评论试图通过误导性的或虚假的陈述欺骗消费者或抬高商家信誉, 因此从庞大的用户评论数据中准确地识别出垃圾评论就显得至关重要, 得到了研究社区和工业界的持续关注。

现有的垃圾信息检测方法大致可以分为 3 类: 基于向量空间的方法^[1-3]、基于语义的方法^[4-7]和基于图神经网络的方法^[8-14]。基于向量空间的方法将垃圾信息检测视为一个文本分类问题, 利用自然语言处理技术从文本中提取特征, 进而应用机器学习算法进行分类。基于向量空间的方法凭借其通用性和灵活性, 在垃圾信息检测领域得到了广泛应用。研究者们不断探索新的文本特征表示和机器学习算法, 持续拓展该方法的适用场景和性能边界。然而, 这类方法仍然存在一些不足。首先, 特征工程在很大程度上依赖领域专家知识, 且对训练数据的质量和数量要求较高。其次, 简单的向量空间表示可能无法充分捕捉文本的语义信息和上下文关系。此外, 面对社交媒体等大规模动态数据, 该类方法在计算效率和实时性方面也面临挑战。为了克服上述局限性, 研究人员探索了基于语义的垃圾信息检测方

基金项目: 中国交通教育研究会 2022—2024 年度教育科学研究重点课题(JT2022ZD043); 南京交通职业技术学院科研项目(JZ2207); 江苏省 2026 年人工智能赋能高等教育新生态专项重点课题(2026XST041)。

收稿日期: 2025-11-25; **修订日期:** 2026-01-23

法,借助于深度学习模型、大语言模型等深入捕捉文本的深层次语义信息,处理复杂的语言表达,进而判断信息性质。随着深度学习技术的发展,卷积神经网络、循环神经网络以及BERT(Bidirectional encoder representations from Transformers)、ALBERT(A lite BERT)等预训练语言模型在文本语义理解方面展现出卓越的表现,极大地推动了垃圾信息检测性能和实用性提升。然而,该类方法主要关注评论文本本身,往往忽略了评论背后的复杂用户行为以及评论与商品之间的关联,无法全面捕捉垃圾评论的多维度特征。因此,在语义理解的基础上,如何充分利用社交网络的结构化信息、跨领域知识等异质信息,并进一步提高模型的鲁棒性、可解释性和计算效率,仍然是亟待解决的问题。

基于图神经网络(Graph neural network, GNN)的方法^[15]为垃圾信息检测开辟了新思路。通过构建包含用户、商品、评论节点的异构图, GNN模型能够灵活地建模不同类型节点间的复杂交互关系,使得垃圾评论的识别可以同时利用文本语义特征与用户行为模式等多源信息。这类方法凭借其强大的关系表示能力,在垃圾信息检测任务上取得了瞩目的成绩。然而,现有的GNN方法在应对类别不平衡问题和刻画多样化节点关系时仍存在局限。特别地,当垃圾评论在图中分布稀疏,或者评论与其他实体之间的关系较少时,模型很难充分学习到垃圾评论的区分性表征,进而影响检测性能。

为了克服上述挑战,本文提出了一种基于图神经网络的垃圾信息检测(GNN-based spam detection, GNNSD)方法。该方法有两个核心组件:(1)设计了一种语义感知的节点增强机制,该机制通过生成语义相关的合成节点来增强少数类别节点的特征表达能力;(2)设计了融合多头关系感知注意力和语义-关系双重注意力的聚合机制,细粒度地建模评论,以及用户和商品之间的复杂交互关系。通过结合这两种机制,GNNSD能更全面地捕捉电商评论系统中的多维异质信息,从而提升对垃圾信息的检测性能。

1 相关工作

在网络环境中信息并非孤立存在,而是与用户以及其他实体存在复杂的交互关系,共同构成了一个异构信息网络。传统的垃圾信息检测方法难以有效地建模和利用这种网络结构信息,而GNN专门用于处理图结构数据,能够直接在图上学习节点表示和边关系。将其引入到垃圾信息检测领域能够更好地捕获多实体间复杂的交互关系,充分挖掘评论中蕴含的语义信息。

首先, GNN在对评论间复杂关联建模方面具有优势。例如, He等^[8]提出的图感知深度融合网络从用户、评论、商品3个维度出发,通过多个子图提取局部和全局特征,采用外积融合技术将低阶和高阶交互信息无缝结合。这种方式不仅保留了关键上下文信息,还提升了模型的泛化能力。

其次, GNN能充分利用丰富的嵌入信息。Cao等^[9]设计的模型创新性地背景知识融入异构图的构建过程中,通过捕获剧情、角色等元素与评论内容之间的语义关联,并巧妙结合文本特征和用户行为数据,实现了性能提升,表明背景知识对垃圾信息识别的重要价值。类似地, Zhang等^[10]专注于用户之间的互动模式,利用图注意力网络揭示协同攻击行为,充分利用用户行为数据,在捕捉隐性垃圾评论方面展现出强大潜力。Andresini等^[11]利用异构图统一地表示用户、商品和评论间的关系,通过图卷积学习丰富的结构特征,在垃圾评论检测任务上取得了良好效果。

最后,在模型结构设计方面,有些学者提出了融合多种注意力机制的方法,以期更全面地捕捉文本语义信息。为克服传统单一注意力机制可能忽略部分重要信息的局限性, Elakkiya等^[12]设计了一种联合注意力机制,综合考虑局部语义和长距离依赖,显著增强了模型对文本语义的理解能力。Liu等^[13]则提出了一个分层注意力网络,通过捕捉文本的多粒度结构信息和句子间交互,进一步提升了垃圾信息检测性能。这些研究表明,精心设计的注意力机制有助于深度学习模型更好地挖掘文本语义特征,提

高垃圾信息检测效果。

综上,从多视角特征融合,到深度特征挖掘,再到复杂关系建模,现有的研究成果已经表明了将GNN应用于垃圾信息检测方法的有效性。然而,评论信息类别不平衡问题的广泛存在使得现有的GNN方法无法充分学习垃圾评论的区分性表征,进而影响了检测方法的性能。

2 垃圾信息检测方法

2.1 问题描述

面向电商平台,将垃圾信息检测方法涉及的主体,包括用户、评论和商品,建模为异构信息网络,形式化地表示为 $G=(V, E, Y)$,其中:节点集合 V 由用户节点集 U 、评论节点集 C 和商品节点集 P 三类节点组成;边集合 E 则包含两类关系:用户与其发表的评论之间的关系集 E_{uc} ,以及评论与其对应商品之间的关系集 E_{cp} ; Y 表示评论节点对应的类别标签集。对于不同类型的节点,其特征表示为:用户节点 $v_u \in U$ 的特征向量为 $\mathbf{x}_u \in \mathbf{R}^m$,评论节点 $v_c \in C$ 的特征向量为 $\mathbf{x}_c \in \mathbf{R}^d$,商品节点 $v_p \in P$ 的特征向量为 $\mathbf{x}_p \in \mathbf{R}^n$ 。其中,评论节点 $v_c \in C$ 是目标节点,具有二元类别标签 $y_c \in \{0, 1\}$, $y_c = 1$ 表示评论 v_c 为垃圾信息, $y_c = 0$ 表示 v_c 为正常评论。

垃圾评论检测的目标是以评论节点为中心,结合评论节点、用户节点和商品节点的特征及其之间的结构化关系,学习式(1)所示的分类函数,即有

$$f: \mathbf{R}^d \times \mathbf{R}^m \times \mathbf{R}^n \rightarrow \{0, 1\} \quad (1)$$

对于每个评论节点 $v_c \in C$,利用其自身特征 $\mathbf{x}_c$ 、与之相连的用户节点特征 $\{\mathbf{x}_u\}_{v_u \in \mathcal{N}_U(v_c)}$ 以及商品节点特征 $\{\mathbf{x}_p\}_{v_p \in \mathcal{N}_P(v_c)}$ 对其类别标签进行预测,有

$$\hat{y}_c = f(\mathbf{x}_c, \{\mathbf{x}_u\}_{v_u \in \mathcal{N}_U(v_c)}, \{\mathbf{x}_p\}_{v_p \in \mathcal{N}_P(v_c)}) \quad (2)$$

式中: $\mathcal{N}_U$ 和 $\mathcal{N}_P$ 分别表示用户节点集合和商品节点集合。

为了学习函数 f ,最小化目标函数为

$$\min_f L(f) = \frac{1}{|C|} \sum_{v_c \in C} l(\hat{y}_c, y_c) \quad (3)$$

式中: $l(\cdot, \cdot)$ 是特定的损失函数, y_c 表示评论节点 v_c 的真实类别标签, $\hat{y}_c$ 表示模型对评论节点的预测标签。该优化目标旨在提高模型对所有评论节点的标签预测准确率。

2.2 方法框架

为了解决电商评论异构图中节点类别不平衡和复杂多类型节点关系建模的问题,本文提出的基于GNN的垃圾信息检测方法框架如图1所示。该方法由4个模块构成:节点增强、特征预变换、基于注意力的信息聚合以及检测器构建。

首先,采用语义感知的节点增强机制,通过个性化PageRank^[16]计算紧密邻居关系,为少数类别节点生成语义相关的合成节点,增强图的结构和特征表达能力。随后,在特征预变换阶段,针对每种节点类型独立设计线性变换,将原始特征映射到统一的隐空间,为后续的聚合操作提供基础。在聚合阶段,模型分别采用多头注意力机制和语义-关系双重注意力机制^[14,17]来捕捉节点间的交互信息。其中,多头注意力机制关注节点间的一般关系模式,而双重注意力机制则同时建模节点的语义特性与关系特性。这两类注意力机制的输出通过拼接融合,从而实现对节点特征的全面表征。为了增强特征的多层表达,模型通过跨层残差连接将不同层的目标节点特征进行拼接以丰富节点表示。最后,使用对比学习优化策略,通过引入正负样本的对比损失增强节点的判别能力,进一步优化嵌入表示。

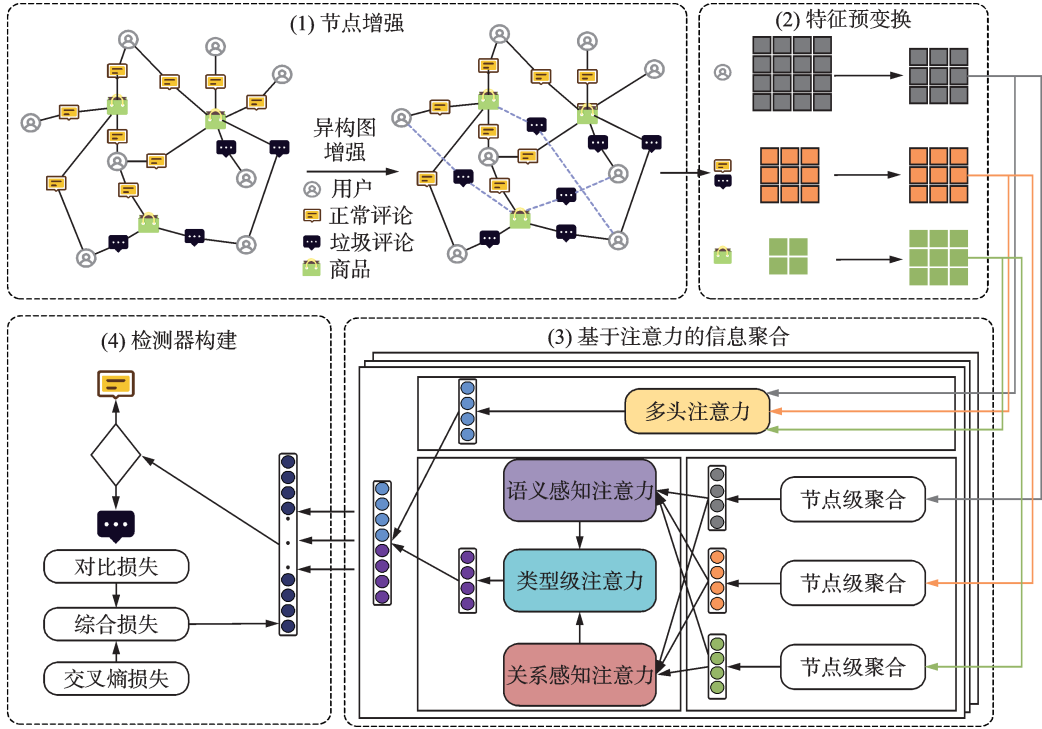


图1 基于图神经网络的垃圾信息检测方法框架

Fig.1 Framework for spam detection method based on graph neural network

2.3 节点增强

电商评论数据通常呈现不平衡分布,少数类节点(如垃圾信息)数量稀少,往往难以被有效学习,导致模型区分能力不足。为缓解这一问题,引入基于语义感知的邻域选择方法,通过生成语义相关的合成节点,增强少数类节点的特征表达能力,同时保持图结构的连贯性和合理性。

为确保生成的合成节点连接到语义相关的邻居,采用基于个性化PageRank^[16]的邻域选择方法。具体而言,对于目标评论节点集 C 和其邻居节点类型集合,构建两者间的一阶对称邻接矩阵 A^R ,计算语义影响力矩阵 Π^R ,该矩阵能捕捉多跳语义信息,而无需显式进行消息传递,表达式为

$$\Pi^R = \alpha(I - (1 - \alpha)\tilde{A}^R)^{-1} \quad (4)$$

式中: $\tilde{A}^R = D^{-1/2} A^R D^{-1/2}$ 为归一化的对称邻接矩阵, D 为度矩阵; $\alpha \in (0, 1]$ 为随机游走的重启概率。

对于少数类别垃圾评论节点,其邻居节点 v_j 的语义影响力度量,表达式为

$$P_j = \sum_{v_i \in S} \Pi_{i,j}^R \quad (5)$$

式中:集合 S 表示所有标签为1的垃圾信息节点,矩阵元素 $\Pi_{i,j}^R$ 表示节点 v_i 对邻居节点 v_j 的语义影响强度。

根据语义影响力得分 P ,选择得分最高的 k 个节点构成候选邻居集合 N^R 。其中 $k = \mu d_{\max}^R$, μ 为控制候选集合大小的超参数, $d_{\max}^R$ 为垃圾信息节点的最大度数。

从候选邻居集合 N^R 中随机抽取若干节点作为合成目标节点的邻居,并根据少数类别的度分布确定邻居数量,以保持图的整体结构一致性。合成节点的特征通过少数类别节点特征的加权平均生成。

2.4 特征预变换

为了消除异构信息网络中用户 U 、评论 C 和商品 P 这 3 种不同类型节点在维度、分布以及编码形式的特征差异,设计了一种节点特征变换方法,为每种节点类型进行独立的线性变换将特征向量映射到统一的隐空间,表达式为

$$\mathbf{x}'_i = \text{ReLU}(\mathbf{W}_{\phi(i)} \mathbf{x}_i) \quad (6)$$

式中: $\mathbf{x}_i$ 表示节点 v_i 的原始属性向量, $\mathbf{x}'_i$ 表示映射到隐空间后的特征表示, $\mathbf{W}_{\phi(i)}$ 表示与节点类型 $\phi(i)$ 对应的可学习变换矩阵。

变换后不同类型节点的特征被映射到一个共享的向量空间,提供了一致的输入,有效地缓解异构网络中节点属性的异质性问题,使模型能够更好地捕捉节点间的交互模式和语义关联。

2.5 基于注意力的信息聚合

为精准捕捉节点间多类型关系的交互特性,设计了一种基于多头关系感知注意力机制的聚合模块。该机制通过显式建模关系类型特性,生成语义丰富且多样性强的节点嵌入,从而充分挖掘节点自身特征与复杂关系模式的潜在价值,增强模型对异构图数据的表达能力。

对于节点 $v_i \in V$,其类型为 $t = \phi(i) \in \{U, C, P\}$,模块定义了与类型 t 对应的查询、键和值这 3 个特征变换矩阵,以适应不同类型节点的特征分布差异。节点 v_i 的查询向量 $\mathbf{q}_i$ 、键向量 $\mathbf{k}_i$ 和值向量 $\mathbf{b}_i$ 分别为

$$\mathbf{q}_i = \mathbf{W}_t^q \mathbf{x}'_i, \quad \mathbf{k}_i = \mathbf{W}_t^k \mathbf{x}'_i, \quad \mathbf{b}_i = \mathbf{W}_t^v \mathbf{x}'_i \quad (7)$$

式中: $\mathbf{x}'_i$ 为节点 v_i 在隐空间的特征表示; $\mathbf{W}_t^q$ 、 $\mathbf{W}_t^k$ 、 $\mathbf{W}_t^v \in \mathbb{R}^{d_{\text{out}} \times d_{\text{in}}}$ 为节点类型 t 的可学习参数矩阵,其中 d_{in} 和 d_{out} 分别表示输入和输出特征的维度。

为了精准捕捉边的语义特性,为每种关系类型 $r \in \mathbf{R} = E_{\text{uc}} \cup E_{\text{cp}}$ 定义了嵌入向量 $\mathbf{e}_r \in \mathbb{R}^{H \times d_h}$,其中 H 表示多头注意力的数量, d_h 为每个注意力头的特征维度。目标节点 v_i 和源节点 v_j 在关系 r 下的注意力得分表达式为

$$\text{Attention}(v_i, v_j, r) = \frac{(\mathbf{q}_i + \mathbf{e}_r) \mathbf{k}_j^T}{\sqrt{d_h}} \quad (8)$$

注意力权重通过 Softmax 函数归一化,以衡量不同邻居节点对目标节点的影响,其中 $\mathcal{N}_i^r$ 为通过关系 r 与目标节点 v_i 相连的邻居节点集合,有

$$\alpha_{ij}^r = \frac{\exp(\text{Attention}(v_i, v_j, r))}{\sum_{v_n \in \mathcal{N}_i^r} \exp(\text{Attention}(v_i, v_n, r))} \quad (9)$$

在得到节点间在不同关系类型下的注意力权重之后,将邻居节点的值向量 $\mathbf{b}_j$ 与对应的注意力权重 α_{ij}^r 相乘加权求和得到目标节点 v_i 的隐空间表示,即

$$\mathbf{z}_i = \sum_{r \in \mathbf{R}} \sum_{v_j \in \mathcal{N}_i^r} \alpha_{ij}^r \mathbf{b}_j \quad (10)$$

这一聚合操作综合了不同语义关系下的节点交互信息,使得节点的隐空间表示能够更全面地反映其在异构图中的语义角色。同时,为了增强模型的优化稳定性并缓解可能出现的梯度消失问题,模块在聚合后的隐空间表示 $\mathbf{z}_i$ 的基础上引入了残差连接,并使用非线性激活函数 GELU 进行变换,得到最终的节点隐藏状态 $\mathbf{h}_i$ 为

$$\mathbf{h}_i = \text{GELU}(\mathbf{z}_i + \mathbf{x}'_i) \quad (11)$$

单一的关系感知注意力机制可能无法充分利用节点的语义信息,限制了节点表示的表达能力。为了进一步提升模型的表示能力,引入结合语义感知与关系感知的双重注意力机制,综合考虑节点的语

义特征与关系类型的复杂依赖,生成更丰富且具判别性的节点表示。

2.5.1 节点级聚合

采用节点级聚合旨在为目标节点生成关系特征表示。为有效建模多关系网络,设计了基于关系类型的聚合方法,确保不同关系类型的信息得到充分表达。具体地,模型将原始异构信息网络按关系类型分解为若干二分子图。每个子图仅保留特定关系类型的边,在子图内聚合节点特征时,为目标节点生成对应关系类型的关系特征。这些关系特征为后续的语义与关系模式建模奠定基础,并通过关系感知建模捕捉节点间复杂的关系类型特性。

为了将目标节点的特征映射到适合聚合的隐空间中,模型引入了与节点类型相关的映射函数。节点 v_i 的隐空间表示为

$$\tilde{h}_i^{(l,k)} = W_{\phi(i)}^{(l,k)} h_i^{(l)} \quad (12)$$

式中: $\tilde{h}_i^{(l,k)}$ 表示节点 v_i 在第 l 层、第 k 个注意力头中的投影后隐空间特征; $W_{\phi(i)}^{(l,k)}$ 为与节点类型 $\phi(i)$ 对应的可学习权重矩阵; $h_i^{(l)}$ 为节点 v_i 在第 l 层的输入特征表示,而初始特征 $h_i^{(0)}$ 则来源于预变换阶段后得到的输出特征 x_i^t 。

在完成隐空间映射后,模型通过注意力机制动态调整邻居节点对目标节点的贡献权重。具体而言,与目标节点 v_i 通过关系 r 相连的邻居节点 v_j 的注意力权重计算表达式为

$$\alpha_{(i,j)}^{(l,k)} = \frac{\exp(\text{ReLU}([\tilde{h}_i^{(l,k)} \| \tilde{h}_j^{(l,k)}] a_n^{(l,k)\text{T}}))}{\sum_{n \in \mathcal{N}_i^r} \exp(\text{ReLU}([\tilde{h}_i^{(l,k)} \| \tilde{h}_n^{(l,k)}] a_n^{(l,k)\text{T}}))} \quad (13)$$

式中: $\alpha_{(i,j)}^{(l,k)}$ 表示节点 v_j 对目标节点 v_i 在第 l 层、第 k 个注意力头中的贡献权重; $a_n^{(l,k)\text{T}}$ 为注意力机制的可学习参数向量。

通过计算注意力权重,模型以加权求和的方式聚合邻居节点的特征,从而生成目标节点的关系特征向量为

$$h_{(i,j)}^{(l,k)} = \sum_{j \in \mathcal{N}_i^r} \alpha_{(i,j)}^{(l,k)} \tilde{h}_j^{(l,k)} \quad (14)$$

2.5.2 类型级聚合

在完成目标节点邻域关系特征的生成后,采用类型级聚合方法整合这些关系特征,以生成最终的节点表示。在电商评论异构网络中,每个目标节点与两个关系类型相关联,分别为用户-评论关系 E_{uc} 和评论-商品关系 E_{cp} ,其特征表示为一组关系向量 $\{h_{i,1}, h_{i,2}\}$ 。

传统的加权平均或求和方法在处理关系向量时,通常假设所有关系类型对目标节点的贡献相等。然而,这一假设在实际应用中往往不成立,因为不同关系类型对目标节点的影响可能存在显著差异。因此,引入了两种互补的注意力机制,以分别捕获基于节点语义的权重和关系类型间的依赖性。

语义感知注意力通过节点语义特征的相似性评估关系向量的影响力。对于第 l 层第 k 个注意力头,节点 v_i 和关系 r 的语义感知注意力权重计算公式为

$$\alpha_{(i,r),s}^{(l,k)} = \frac{\exp(\text{ReLU}([\tilde{h}_i^{(l,k)} \| h_{(i,r)}^{(l,k)}] a_s^{(l,k)\text{T}}))}{\sum_{n \in R_i} \exp(\text{ReLU}([\tilde{h}_i^{(l,k)} \| h_{(i,r)}^{(l,k)}] a_n^{(l,k)\text{T}}))} \quad (15)$$

式中: $\alpha_{(i,r),s}^{(l,k)}$ 表示节点 v_i 与关系 r 的语义感知注意力权重; $a_s^{(l,k)\text{T}}$ 为语义注意力的可学习参数向量; R_i 表示节点 v_i 的关联关系集合。

关系感知注意力旨在从关系类型的上下文中捕获依赖信息。对于关系 r 的关系感知注意力权重,计算公式为

$$\alpha_{(i,r),r}^{(l,k)} = \frac{\exp(\text{ReLU}(\mathbf{e}_r^{(l)} \mathbf{a}_r^{(l,k)\top}))}{\sum_{n \in R_i} \exp(\text{ReLU}(\mathbf{e}_r^{(l)} \mathbf{a}_n^{(l,k)\top}))} \quad (16)$$

式中: $\alpha_{(i,r),r}^{(l,k)}$ 表示节点 v_i 和关系 r 的关系感知注意力权重; $\mathbf{e}_r^{(l)}$ 表示关系 r 的嵌入表示; $\mathbf{a}_r^{(l,k)\top}$ 为关系注意力的可学习参数向量。

为了综合语义感知和关系感知注意力,提出类型级注意力的融合方法,计算公式为

$$\alpha_{(i,r)}^{(l,k)} = \lambda_k \alpha_{(i,r),s}^{(l,k)} + (1 - \lambda_k) \alpha_{(i,r),r}^{(l,k)} \quad (17)$$

式中 λ_k 为平衡系数,可以通过模型训练动态调整。

在计算出类型级注意力分数后,模型以加权求和的方式对关系向量进行聚合,表达式为

$$\tilde{\mathbf{h}}_i^{(l+1)} = \left\| \sum_{k=1}^K \alpha_{(i,r)}^{(l,k)} \mathbf{h}_{(i,r)}^{(l,k)} \right\| \quad (18)$$

$$\mathbf{h}_i^{(l+1)} = \text{Norm}(\sigma(\tilde{\mathbf{h}}_i^{(l+1)} + \mathbf{h}_i^{(l)})) \quad (19)$$

式中: $\left\| \sum_{k=1}^K \right\|$ 表示多头注意力的拼接操作, $\sigma(\cdot)$ 为 Sigmoid 激活函数, $\text{Norm}(\cdot)$ 表示 L_2 正则化。

通过类型级聚合方法,模型有效地结合了节点语义信息和关系依赖性,从而生成了更具代表性的节点嵌入表示。为整合多头关系感知注意力机制和语义-关系双重感知注意力机制所生成的嵌入,进一步地将多头关系感知嵌入 $\mathbf{h}_i^m$ 与双重感知注意力机制嵌入 $\mathbf{h}_i^b$ 拼接后,经线性变换与 GELU 激活函数映射为最终节点表示,有

$$\mathbf{h}_i^c = \text{GELU}(\mathbf{W}_c \cdot [\mathbf{h}_i^m \| \mathbf{h}_i^b]) \quad (20)$$

式中 $\mathbf{W}_c$ 为用于线性变换的可学习权重矩阵。

2.6 检测器构建

为充分利用不同网络层的信息,检测器模型采用跨层残差连接策略,将每一层生成的目标节点嵌入 $\mathbf{h}_i^c$ 进行拼接以增强节点嵌入的多样性和表达能力,表达式为

$$\mathbf{h}_i^t = \text{concat}(\mathbf{h}_i^{c(0)}, \mathbf{h}_i^{c(1)}, \dots, \mathbf{h}_i^{c(L)}) \quad (21)$$

式中: L 为模型的总层数, concat 表示特征向量的拼接操作。通过残差连接确保梯度在深层网络中的有效传递,有效缓解梯度消失问题,从而确保模型的训练稳定。

模型的损失函数考虑交叉熵损失和对比损失,旨在增强节点嵌入的判别能力,进而优化节点检测性能。交叉熵损失用于监督节点分类任务,通过最小化真实类别分布与预测类别分布之间的差异,使模型能够准确捕捉节点的类别特征,计算公式为

$$\mathcal{L}_{\text{CE}} = -\frac{1}{N} \sum_{i=1}^N \sum_{c=1}^2 y_{i,c} \lg \hat{y}_{i,c} \quad (22)$$

式中: N 为目标节点的数量, $y_{i,c}$ 表示节点 v_i 的真实类别标签, $\hat{y}_{i,c}$ 为模型预测的概率。通过这一损失函数,模型能够捕捉节点的类别特征,提升分类精度。

对比损失通过优化目标节点嵌入与正负样本嵌入之间的相似性,增强节点表示的判别能力。具体而言,对比损失函数鼓励目标节点的嵌入与正样本节点的嵌入在余弦相似度空间中更接近,而与负样本节点的嵌入在余弦相似度空间中更远离。通过这种方式,对比损失有效地提升了节点嵌入在语义空间中的区分能力。其损失函数形式为

$$\mathcal{L}_{\text{Contrastive}} = -\frac{1}{|T|} \sum_{i \in T} \lg \sigma(h_i^t g^+ - h_i^t g^-) \quad (23)$$

式中: T 表示目标节点集合,捕捉全局正样本特征; $\mathbf{h}_i^t$ 表示节点 v_i 的最终嵌入表示, $\mathbf{h}_j^t$ 表示与目标节点 v_i

进行对比的其他节点的最终嵌入表示; $g^+ = \frac{1}{|T|} \sum_{j \in T} h'_j$ 和 $g^- = \frac{1}{|T|} \sum_{j \in T'} h'_j$ 分别为正样本和负样本的全局表示, T' 为目标节点集合的乱序版本用于模拟负样本全局特征。

综合交叉熵损失和对比损失的函数表达为

$$\mathcal{L} = \mathcal{L}_{CE} + \lambda \mathcal{L}_{Contrastive} \quad (24)$$

式中 λ 为权重参数,用于平衡分类任务与对比学习的贡献。

3 实 验

3.1 实验设定

采用来自Yelp平台的基准数据集Yelp-Hotel和Yelp-Restaurant^[18]开展实验验证方法有效性。这两个数据集包含了用户提交的文本评论,Yelp-Hotel包含来自72家酒店的评论,而Yelp-Restaurant则包含来自129家餐馆的评论。每条评论都与对应的商户和评论者关联。数据集还包含丰富的元数据,包括评论者信息(如好友数量、评论数量)以及商户属性(如产品类别、价格范围和评分)。数据集呈现明显的类别不平衡特征,其中垃圾评论(spam)的数量显著少于非垃圾评论(non-spam),Yelp-Hotel和Yelp-Restaurant中虚假评论仅占约14%。为保证实验结果的可靠性,采用了60%-20%-20%的随机划分策略,将数据分为训练集、验证集和测试集。

考虑到数据存在严重的类别不平衡问题,选择了4个适合评估不平衡分类问题的指标:AUC、Micro- F_1 、Macro- F_1 和G-mean,避免造成类别不平衡引起的评估偏差,解释如下。

AUC:衡量模型区分不同类别的能力,其值不受类别分布影响,其中 M 和 N 分别表示正样本和负样本的数量, rank_i 表示第 i 个正样本的排名。

$$\text{AUC} = \frac{\sum_{i \in \text{positiveClass}} \text{rank}_i - \frac{M \times (1 + M)}{2}}{M \times N} \quad (25)$$

式中positiveClass表示标注为虚假评论的样本集。

Micro- F_1 :将所有类别的样本视为整体计算 F_1 值,适合评估类别不平衡数据的整体性能,其中Pre和Recall分别表示精确率和召回率。

$$\text{Micro-}F_1 = \frac{2 \times \text{Pre} \times \text{Recall}}{\text{Pre} + \text{Recall}} \quad (26)$$

Macro- F_1 :对每个类别单独计算 F_1 后取平均,确保少数类在评估中得到同等重视,其中 S 为类别数量, Pre_i 和 Recall_i 分别表示第 i 个类别上的精确率和召回率。

$$\text{Macro-}F_1 = \frac{1}{S} \sum_{i=1}^S \frac{2 \times \text{Pre}_i \times \text{Recall}_i}{\text{Pre}_i + \text{Recall}_i} \quad (27)$$

G-mean:通过计算各类别准确率的几何平均值,特别关注模型在少数类上的表现,其中 Accuracy_i 表示第 i 个类别上的精度。

$$\text{G-mean} = \sqrt[S]{\prod_{i=1}^S \text{Accuracy}_i} \quad (28)$$

同时,选择3个具有代表性的基于语义的方法和5个基于图神经网络的方法作为对比方法。选取的方法如下。

TextCNN^[19]是一种基于卷积神经网络的文本分类模型,通过卷积和池化操作提取文本局部特征,在文本分类任务中表现优异。

ALBERT-BiLSTM^[20]融合了ALBERT的语义表示和BiLSTM捕捉上下文信息的优势,并通过注

意力机制突出关键词,增强了短文本分类能力。

EUPHORIA^[21]是一个多视角垃圾评论检测框架,结合评论内容(Word2Vec和BERT特征)和评论者行为视角,通过多输入神经网络进行特征融合,大幅提升分类性能。

GraphSAGE^[22]通过采样和聚合邻居节点信息生成节点嵌入表示,是一种经典的可扩展图神经网络模型,适用于大规模图数据处理。

CARE-GNN^[23]引入标签感知相似度和强化学习优化的邻居选择策略,增强了对伪装欺诈者的检测能力,是一种专门为欺诈检测设计的图神经网络。

PC-GNN^[24]通过标签平衡采样和邻居选择机制缓解类不平衡问题,增强了少数类欺诈节点的表示学习,提高了检测准确率。

GLORIA^[11]是一种基于图卷积网络的垃圾评论检测方法。它利用异构图统一表示用户、商品和评论间的关系,通过图卷积学习丰富的结构特征,在垃圾评论检测任务上取得了良好效果。

IN-GFD^[25]同样是一种基于图卷积网络的垃圾评论检测方法,它通过过采样垃圾评论节点来使其与正常节点平衡,并引入了一个边损失函数来学习新的边关系,进而消除类别不平衡的问题的影响。

3.2 性能评价及对比

表1描述了本文提出的方法整体性能,并与选取的对比方法进行了对比分析。从整体性能来看,GNNSD在两个数据集的各项评价指标上均展现出较好的性能。以Yelp-Hotel数据集为例,GNNSD的AUC、Micro- F_1 、Macro- F_1 和G-mean分别达到0.921 9、0.906 8、0.753 3和0.734 3。与表现最好的基于语义的方法EUPHORIA相比,分别提升了6.28%、4.71%、11.8%和12.25%,主要原因在于GNNSD借助于图结构捕获评论、用户和商品之间的关联信息的方法,能弥补仅考虑评论内容的语义信息的不足,大幅提升性能。其他基于图神经网络的方法也有类似的结论,这也验证了引入图表示学习范式在垃圾信息检测领域的优越性。与这些基于图神经网络的方法相比,GNNSD在AUC、Micro- F_1 和Macro- F_1 指标上都获得了最佳表现,这表明GNNSD在捕获异构图中丰富语义和结构信息、挖掘不同类型节点复杂交互关系方面较为优异。

表1 模型性能与对比结果

Table 1 Performance of the proposed model and comparison models

模型	Yelp-Hotel				Yelp-Restaurant			
	AUC	Micro- F_1	Macro- F_1	G-mean	AUC	Micro- F_1	Macro- F_1	G-mean
TextCNN	0.534 5	0.865 8	0.534 6	0.231 8	0.587 2	0.761 7	0.595 3	0.490 9
ALBERT-BiLSTM	0.570 3	0.870 8	0.589 0	0.397 4	0.642 7	0.799 7	0.663 7	0.566 2
EUPHORIA	0.859 1	0.859 7	0.635 3	0.511 8	0.908 5	0.900 3	0.717 8	0.644 1
GraphSAGE	0.810 3	0.793 1	0.668 1	0.741 0	0.837 0	0.818 1	0.679 7	0.782 3
CARE-GNN	0.840 8	0.762 5	0.661 8	0.777 8	0.877 9	0.786 5	0.666 1	0.816 9
PC-GNN	0.861 6	0.805 2	0.683 1	0.756 6	0.889 3	0.860 8	0.715 7	0.809 9
GLORIA	0.887 6	0.868 2	0.747 3	0.742 6	0.925 8	0.896 2	0.762 3	0.834 2
IN-GFD	0.901 2	0.889 3	0.752 9	0.747 5	0.948 1	0.913 8	0.852 1	0.839 9
GNNSD	0.921 9	0.906 8	0.753 3	0.734 3	0.963 0	0.941 9	0.861 5	0.859 5

进一步地,分析不同图神经网络方法的表现差异,可以发现图的异质性建模在提升检测精度方面至关重要。与同构图方法GraphSAGE相比,CARE-GNN、PC-GNN和GLORIA在异构图上设计了特定的邻居聚合和关系感知机制,充分利用了不同类型节点和关系的特征,使得检测性能得到显著改善。

这表明现实场景中的垃圾信息网络往往呈现出较强的异质性,异构图神经网络能更好地适应这一特点。尽管如此,这些异构图方法的表现仍不及 GNNSD。究其原因,GNNSD 在异构图建模的基础上,进一步融合了全面的语义增强策略(如节点数据增广)和细粒度关系感知机制(如语义-关系双重注意力),从而能够更加准确地刻画复杂、细腻节点交互模式,使得检测性能得到进一步提升。

此外,GNNSD 在对比学习中充分融合了节点的语义相似度和结构相似度信息,通过优化节点的结构表示,提升了节点嵌入的判别能力和抗干扰性。相比之下,尽管 EUPHORIA 整合了文本语义和用户行为特征,但在显式建模图结构信息方面仍存在不足。GNNSD 采用多视角学习范式,不仅有效融合了异质信息,还充分挖掘了图的结构特性,从而在模型的表达能力上展现出独特的优势。实验结果表明,这种多视角对比学习方法能够显著增强垃圾信息检测模型的泛化性能。

需要指出的是,相比 Yelp-Restaurant 数据集,所有模型在 Yelp-Hotel 数据集上的性能均有所下降。这是因为 Yelp-Restaurant 对应的异构图在节点数量和类型上都更加丰富,有利于挖掘复杂的交互模式,而 Yelp-Hotel 的异构图则相对简单,缺乏一些有判别力的结构特征。即便如此,GNNSD 在 Yelp-Hotel 上的表现仍较其他方法高出约 3%,彰显了其在稀疏场景下的方法鲁棒性,使得其更加实用。

3.3 消融实验

为评估 GNNSD 各个关键组件对模型性能的影响,设计了一组消融实验。通过逐一移除或替换模型的特定模块,构建了以下 4 个模型变体:(1)w/o NA(without Node augmentation)表示移除节点数据增强策略;(2)w/o RHA(without Relation-aware multi-head attention)表示移除多头关系感知注意力机制;(3)w/o DASE(without Dual attention for semantics and relations)表示移除语义感知与关系感知的双重注意力机制;(4)w/o CL(without Contrastive learning)表示移除对比学习模块。

通过比较这些变体与 GNNSD 的性能差异,可以分析每个组件的贡献。表 2 展示了在两个数据集上的消融实验结果。可以看出,GNNSD 在大多数评估指标上均取得了最优表现,验证了模型设计的合理性和有效性。具体地,多头关系感知注意力和语义-关系双重注意力的移除对模型性能影响最为显著,尤其体现在对异质信息建模能力敏感的指标,如 AUC、Macro- F_1 和 G-mean 上。以 Yelp-Restaurant 为例,w/o RHA 使 AUC、Macro- F_1 和 G-mean 分别下降了 3%、7.86% 和 7.21%,w/o DASE 的降幅也达 2.29%、5.3% 和 6.7%。这表明捕捉异构图中复杂、多样的节点关系交互模式,并融合节点语义与结构特征,是 GNNSD 实现高效异质表征学习、提升垃圾信息检测性能的关键。

表 2 不同组件对于模型整体性能的贡献
Table 2 Contributions of each component to the performance of GNNSD

模型	Yelp-Hotel				Yelp-Restaurant			
	AUC	Micro- F_1	Macro- F_1	G-mean	AUC	Micro- F_1	Macro- F_1	G-mean
w/o NA	0.914 6	0.883 7	0.727 5	0.664 8	0.944 2	0.934 7	0.829 5	0.828 5
w/o RHA	0.882 9	0.896 5	0.697 7	0.627 3	0.933 0	0.912 1	0.782 9	0.787 4
w/o DASE	0.894 8	0.883 7	0.709 6	0.699 5	0.940 1	0.927 2	0.808 5	0.792 5
w/o CL	0.901 5	0.893 9	0.696 3	0.632 3	0.952 9	0.945 5	0.850 6	0.823 2
GNNSD	0.921 9	0.906 8	0.753 3	0.734 3	0.963 0	0.941 9	0.861 5	0.859 5

节点数据增强和对比学习的贡献相对次要,但也展现出一定的积极效果。特别是在规模较小的 Yelp-Hotel 数据集上,w/o NS 的 Macro- F_1 、G-mean 比 GNNSD 分别下降了 2.58%、6.95%,w/o CL 则下降了 5.7% 和 10.2%;而在数据规模更大的 Yelp-Restaurant 数据集上,NA 和 CL 的贡献相对较小。当图的规模和复杂程度较低时,NA 和 CL 可以在一定程度上弥补数据稀疏性,通过对稀少类别节点的语义

感知增强和对比学习范式,有效提高分类的均衡性。而在大规模异构图上,模型能从丰富的结构化信息中学习到更加鲁棒的节点表征,降低对额外增强机制的依赖。

值得注意的是,尽管个别变体在特定指标上的性能降幅看似较小,如去除 NA 在 Yelp-Restaurant 的 AUC 仅降低 1.88%,但纵观全局,所有变体的整体表现仍然显著优于之前实验中的各基线方法,这证明了 GNNSD 的鲁棒性和泛化能力。此外,消融结果还显示,不同组件的移除对 $\text{Micro-}F_1$ 的影响相对有限,这可能是由于 $\text{Micro-}F_1$ 作为一个宏观指标,对不同类别的样本一视同仁,未能充分反映少数类优化的细微变化。

4 结束语

为了应对垃圾信息检测所面临的类别不平衡与复杂关系建模等挑战,本文提出了一种异构图神经网络模型 GNNSD。该模型引入语义感知节点增强方法,利用丰富的语义信息为少数类别节点生成优质合成邻居,增强其特征表达力,并保持图的整体结构一致性减轻类别不平衡对模型训练的负面影响。同时,引入多头关系感知注意力与语义-关系双重注意力机制,借助于两类注意力机制的互补融合使得节点嵌入能够更全面地表征其语义属性与结构角色。在关键指标上超过现有主流方法证明 GNNSD 在处理复杂异构图结构方面的有效性与鲁棒性。当然,GNNSD 选用的数据视角还不够丰富,仅包括文本内容和用户行为,而恶意信息通常表现出多样化特征,因此下一步融合更加丰富的数据视角并且解决不同视图数据的异质性给建模带来的问题,将有助于提升恶意评论检测准确性、增强模型对恶意模式的理解和判别能力。

参考文献:

- [1] 陈龙, 梁意文, 谭成予. 基于自适应性分类器的垃圾邮件检测[J]. 计算机工程, 2018, 44(5): 194-200.
CHEN Long, LIANG Yiwen, TAN Chengyu. Spam detection based on adaptive classifier[J]. Computer Engineering, 2018, 44(5): 194-200.
- [2] EZPELETA E, ZURUTUZA U, GÓMEZ HIDALGO J M. Does sentiment analysis help in Bayesian spam filtering? [C]// Proceedings of Hybrid Artificial Intelligent Systems. Cham: Springer, 2016: 79-90.
- [3] AIYAR S, SHETTY N P. N-gram assisted Youtube spam comment detection[J]. Procedia Computer Science, 2018, 132: 174-182.
- [4] SAIDANI N, ADI K, ALLILI M S. A semantic-based classification approach for an enhanced spam detection[J]. Computers & Security, 2020, 94: 101716.
- [5] 俞荧妹, 禹素萍, 许武军, 等. 基于深度学习的垃圾邮件检测[J]. 计算机科学与应用, 2023(4): 764-772.
YU Yingmei, YU Suping, XU Wujun, et al. Spam detection based on deep learning[J]. Computer Science and Application, 2023(4): 764-772.
- [6] SHAABAN M A, HASSAN Y F, GUIRGUIS S K. Deep convolutional forest: A dynamic deep ensemble approach for spam detection in text[J]. Complex & Intelligent Systems, 2022, 8(6): 4897-4909.
- [7] SINGHAL R, KASHEF R. A weighted stacking ensemble model with sampling for fake reviews detection[J]. IEEE Transactions on Computational Social Systems, 2024, 11(2): 2578-2594.
- [8] HE L, XU G, JAMEEL S, et al. Graph-aware deep fusion networks for online spam review detection[J]. IEEE Transactions on Computational Social Systems, 2023, 10(5): 2557-2565.
- [9] CAO H, LI H, HE Y, et al. GCMK: Detecting spam movie review based on graph convolutional network embedding movie background knowledge[C]//Proceedings of Artificial Neural Networks and Machine Learning—ICANN 2022. Cham: Springer, 2022: 494-505.
- [10] ZHANG L, SONG X, ZHAO X, et al. GAIM: Graph-aware feature interactional model for spam movie review detection[C]//Proceedings of 2022 26th International Conference on Pattern Recognition (ICPR). Montreal, QC, Canada: IEEE, 2022: 621-628.

- [11] ANDRESINI G, APPICE A, GASBARRO R, et al. GLORIA: A graph convolutional network-based approach for review spam detection[M]//Discovery Science. Cham: Springer Nature Switzerland, 2023: 111-125.
- [12] ELAKKIYA E, SELVAKUMAR S, LEELA VELUSAMY R. TextSpamDetector: Textual content based deep learning framework for social spam detection using conjoint attention mechanism[J]. Journal of Ambient Intelligence and Humanized Computing, 2021, 12(10): 9287-9302.
- [13] LIU Y, WANG L, SHI T, et al. Detection of spam reviews through a hierarchical attention architecture with N-gram CNN and Bi-LSTM[J]. Information Systems, 2022, 103: 101865.
- [14] LI X, ZHANG P, MA R, et al. Integrating heterogeneous graph attention network with label propagation for detecting spammer groups on E-commerce platforms[J]. ACM Transactions on Knowledge Discovery from Data, 2025, 19(8): 1-28.
- [15] 杨永鹏, 杨震, 杨真真. 基于时序分解和注意力图神经网络的交通预测[J]. 数据采集与处理, 2025, 40(2): 417-430.
YANG Yongpeng, YANG Zhen, YANG Zhenzhen. Time-series decomposition and attention graph neural network based traffic forecasting[J]. Journal of Data Acquisition and Processing, 2025, 40(2): 417-430.
- [16] 袁野, 陈明, 吴安彪, 等. 基于个性化PageRank和对比学习的图异常检测模型[J]. 计算机科学, 2025, 52(2): 80-90.
YUAN Ye, CHEN Ming, WU Anbiao, et al. Graph anomaly detection model based on personalized PageRank and contrastive learning[J]. Computer Science, 2025, 52(2): 80-90.
- [17] 韩普, 刘森嶺, 陈文祺. 基于多尺度注意力和图神经网络的多模态医学实体识别研究[J]. 数据采集与处理, 2025, 40(4): 922-933.
HAN Pu, LIU Senling, CHEN Wenqi. Multi-modal medical entity recognition based on multi-scale attention and graph neural networks[J]. Journal of Data Acquisition and Processing, 2025, 40(4): 922-933.
- [18] MUKHERJEE A, VENKATARAMAN V, LIU B, et al. What Yelp fake review filter might be doing?[C]//Proceedings of the International AAAI Conference on Web and Social Media. [S.l.]: AAAI, 2013.
- [19] WANG Z, YU D, LI Q, et al. SR-HGN: Semantic-and relation-aware heterogeneous graph neural network[J]. Expert Systems with Applications, 2023, 224: 119982.
- [20] XU G, ZHOU D, LIU J. Social network spam detection based on ALBERT and combination of Bi-LSTM with self-attention [J]. Security and Communication Networks, 2021, 2021(1): 5567991.
- [21] ANDRESINI G, IOVINE A, GASBARRO R, et al. EUPHORIA: A neural multi-view approach to combine content and behavioral features in review spam detection[J]. Journal of Computational Mathematics and Data Science, 2022, 3: 100036.
- [22] HAMILTON W L, YING R, LESKOVEC J. Inductive representation learning on large graphs[C]//Proceedings of the 31st International Conference on Neural Information Processing Systems. Long Beach, California, USA: ACM, 2017: 1025-1035.
- [23] DOU Y, LIU Z, SUN L, et al. Enhancing graph neural network-based fraud detectors against camouflaged fraudsters[C]//Proceedings of the 29th ACM International Conference on Information & Knowledge Management. Virtual Event, Ireland: ACM, 2020: 315-324.
- [24] LIU Z, DOU Y, YU P S, et al. Alleviating the inconsistency problem of applying graph neural network to fraud detection[C]//Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval. [S.l.]: ACM, 2020: 1569-1572.
- [25] YU H, LIU W, ZHU N, et al. IN-GFD: An interpretable graph fraud detection model for spam reviews[J]. IEEE Transactions on Artificial Intelligence, 2024, 5(10): 5325-5339.

作者简介:



张鸽(1978-), 通信作者, 男, 副教授, 研究方向: 计算机网络, E-mail: 4033830@qq.com。



郑莹(1975-), 女, 副教授, 研究方向: 现代通信技术, E-mail: 279265445@qq.com。



宋万里(1981-), 男, 副教授, 研究方向: 数据分析、智慧教育等, E-mail: swl@njxzc.edu.cn。

(编辑: 张黄群)