



研究与开发

基于大语言模型与事件演化图谱的诈骗电话识别方法

胡程忆

(卓望信息技术(北京)有限公司, 北京 100072)

摘 要: 针对电信诈骗话术快速演化、传统方法依赖大量标注样本以及基于大模型的端到端识别存在幻觉风险等问题, 提出一种基于大语言模型与事件演化图谱的诈骗电话识别方法。该方法引入单类学习思想, 在无需负样本的条件下, 利用大语言模型的零样本能力, 将主叫话术抽象为标准化的动宾短语事件演化链, 通过语义相似度融合构建可增量扩展的事件演化图谱, 并采用关键事件节点匹配与连通性检测进行风险判别。实验结果表明, 该方法的 $F1$ 值在 2 000 条样本规模下达 84.84%, 在 8 000 条样本规模下可提升至 91.46%, 性能随样本规模增长呈现一定的缩放律特征。该方法为识别持续演化的复杂诈骗行为提供了精准且具有扩展性的技术路径。

关键词: 大语言模型; 事件演化图谱; 诈骗电话

中图分类号: TP391

文献标志码: A

doi: 10.11959/j.issn.1000-0801.DXKX250564

Fraud call identification based on large language models and event evolution graphs

Hu Chengyi

Aspire Information Technology (Beijing) Co., Ltd., Beijing 100072, China

Abstract: In response to the rapid evolution of fraud call scripts, the reliance on a large number of labeled samples in traditional methods, and the potential hallucination risks in end-to-end recognition based on large models, a fraud call identification method based on large language models and event evolution graphs was proposed. By introducing the concept of one-class learning, the zero-shot capability of large language models was leveraged to abstract caller speech into standardized verb-object phrase event evolution chains without the need for negative samples. An incrementally expandable event evolution graph was constructed through semantic similarity fusion, and risk discrimination was performed using key event node matching and connectivity detection. Experimental results show that the proposed method achieves an $F1$ score of 84.84% on a sample scale of 2 000 instances, which is further improved to 91.46% on a scale of 8 000 instances, with performance exhibiting certain scaling law characteristics as the sample size increases. The proposed method provides a precise and scalable technical approach for identifying complex and continuously evolving fraudulent behaviors.

Key words: large language model, event evolutionary graph, fraud call



0 引言

电信网络诈骗已成为当前社会面临的重大公共安全问题,其犯罪形式呈现出跨境化、集团化和技术化的特征,对社会经济秩序和公民财产安全构成严重威胁。2020—2023年,以电信网络诈骗为主的诈骗案件在全国公安机关刑事案件立案数中占比40%左右(警情数占比60%),跃居刑事案件立案总数的第一位^[1]。诈骗犯罪不仅波及面广,而且具有线上线下交融、集团化运作、技术性强等特点,造成的经济损失和社会影响极为严重。近年来,国内反诈力度持续加大,诈骗窝点加速向境外转移,不法分子形成以“工业园区”“科技园区”为掩护的跨境犯罪集团,进一步加大了打击治理工作的难度^[2]。从技术层面来看,电信诈骗的手段不断翻新,犯罪分子利用区块链、远程操控、人工智能、虚拟货币等新业态、新技术,不断翻新诈骗手法,升级作案工具,变换洗钱方式,极大增加了侦查打击和追赃挽损的难度^[2]。例如,诈骗分子通过伪造身份、使用变声软件等手段,使被害人难以辨别真伪。此外,诈骗集团内部组织严密,分工明确,形成了从个人信息获取、技术支撑到资金转移的完整产业链,进一步加大了案件侦破和证据固定的难度。

在此背景下,如何高效识别和防范电信诈骗成为亟待解决的问题。尽管近年来诈骗手段不断升级,呈现出技术化、智能化的发展趋势,但电话通信依然是诈骗分子实施犯罪的主要渠道。值得注意的是,无论诈骗手法如何迭代演变,基于通话内容的语义识别始终是最具实效性的技术路径之一。诈骗分子通过电话直接建立与受害者的“一对一”接触,其话术体系往往呈现出显著的语言特征:高频使用威胁性话术制造心理压迫,运用诱导性话术虚构利益诱惑,或通过伪造权威身份获取信任。这些特定的语言模式不仅构成了

诈骗行为的核心要素,更为基于人工智能的内容识别技术提供了关键的特征维度。因此,深入挖掘诈骗电话的语义特征和行为模式,构建智能化的识别模型,对于提升反诈工作的精准度和时效性具有重要的实践价值。

1 相关研究

基于结构化行为特征建模的方法主要通过挖掘通话记录、时空属性及关系结构,将诈骗行为抽象为可计算的特征或图结构,以实现对诈骗用户和号码的自动识别与预警。例如,张照星等^[3]通过连通图法、统计分析和数据挖掘分类算法,基于运营商国际漫游话单数据,拓展电信诈骗嫌疑号码的识别方法,以提高预警和封停效率。李洋等^[4]提出了一种基于攻击经济学分析诈骗电话时空特征并结合机器学习分类的方法,有效检测移动虚拟运营商中的诈骗用户,显著降低了诈骗比例。近年来,学术界对结构化数据建模开始采用图结构融合深度学习的方法进行识别。例如,刘晨迪等^[5]提出了一种融合图结构和多通道注意力机制的诈骗电话识别模型,通过建模用户通话关系图并结合多种神经网络特征提取方法,显著提升了诈骗电话的识别准确率。不过,该方法高度依赖人工特征设计和历史结构模式,对诈骗话术的动态变化和演化路径的适应性有限,且在小样本或跨场景条件下泛化能力不足。

基于语义建模的方法主要从通话文本语义出发,利用深度学习或知识表示技术刻画诈骗话术的语言模式与语义结构,以提升对诈骗内容的理解与判别能力。例如,周俊杰等^[6]提出了一种结合位置编码、双向门控循环单元、卷积神经网络和注意力机制的融合模型,用于提升诈骗电话文本分类的准确性和对关键信息的捕捉能力。司海平等^[7]提出了一种结合提示模板与对比学习的小样本分类方法,通过增强数据表达和优化特征学习,提升了电信诈骗文本的细粒度分类准确率。

基于语义的知识图谱技术也开始发展,例如,邓诗琦等^[8]提出了一种面向智能应用的领域本体构建方法——应用驱动循环法,以智能应用需求为核心驱动跨领域知识融合建模,采用“需求+构建+评估”的循环结构,并以反电信诈骗领域为例,验证了该方法在提升语义理解与应用支撑能力方面的有效性。刘卓娴等^[9]提出了一种基于小样本微调UIE框架的电信网络诈骗警情要素抽取方法,并结合知识图谱构建与频繁子图挖掘算法,有效提升了公安机关对电信网络诈骗警情的智能化分析和犯罪模式发现能力。然而,此类方法普遍依赖充足标注数据和固定语义模式,难以刻画诈骗行为的动态演进过程,这也促使本文进一步引入事件演化图谱与大语言模型进行统一建模。

基于大模型建模的方法主要依托大语言模型的强语义理解与推理能力,直接对诈骗文本或多模态信息进行端到端建模,从而提升对复杂话术和诈骗手段的识别能力。例如,李衡峰等^[10]提出了一种基于大模型的电信网络诈骗文本分类方法,该方法利用句向量聚类 and 自动提示词生成优化技术,仅需少量样本即可高效识别新型诈骗手段,显著提升了预警能力。然而,Shen等^[11]的研究表明,尽管大模型在理解对话上下文方面优于传统方法,但仍面临数据集偏差、召回率低和模型幻觉等挑战。对此,Singh等^[12]提出了一种基于检索增强生成技术的大语言模型系统,用于实时检测电话诈骗。该系统通过动态更新策略文档来验证通话内容合规性及呼叫者身份,无须重新训练模型,并在其特定领域测试中取得了高准确率(97.98%)和高F1值(97.44%)。近年来,慢思考和多模态大模型技术取得了较大进展,Ma等^[13]提出了首个开源的音频-文本慢思考数据集TeleAntiFraud-28k,通过多策略数据生成和慢思考标注机制,为电信诈骗检测提供了多模态分析基础。在大语言模型结合知识图谱方面,裴炳森

等^[14]基于ChatGPT和知识图谱提出了一种电信诈骗案件类型影响力评估方法,通过结构化案件文本、量化案发时间、涉案金额和涉案人数等关键因素,计算影响因子,以科学评估不同诈骗类型的危害程度与演变趋势,为优化反诈策略提供数据支持。现有方法以大模型直接判别为主,存在可解释性不足、幻觉风险等问题,且难以显式刻画诈骗行为的演化结构。因此,本文引入事件演化图谱对大模型能力进行约束与增强,以弥补上述不足。

基于事件演化图谱建模的方法通过将诈骗行为抽象为事件及其演化关系,利用图结构刻画诈骗过程的时序逻辑与因果模式,从而支持对复杂诈骗行为的结构化理解与分析。Guan等^[15]认为,相较于以实体为中心的传统知识图谱(knowledge graph, KG),以事件为中心的知识图谱更擅长表达动态事件知识。事件演化图^[16]和事件逻辑图^[17]则更侧重于抽象事件模式与演化逻辑,多在模式层进行构建和推理。周胜利等^[18]利用知识图谱和事理图谱技术,实证分析了生成式人工智能如何驱动电信网络诈骗风险的演化,识别出3种风险演化模式,为新型诈骗的防治提供了理论依据。斯彬洲等^[19]提出了一种基于大语言模型的事件抽取与融合方法,用于自动构建电信诈骗事件演化链条,并量化关键风险因素,以支撑诈骗预警和防范治理。然而,现有研究多集中于事件模式分析或风险研判层面,事件获取与图谱构建成本较高,且与具体电话识别任务结合不够紧密,难以直接用于实时识别场景。

综合来看,现有研究在结构化行为特征、语义建模、大模型判别及事件图谱分析等方面取得了积极进展,但仍存在若干问题:一是高度依赖人工特征或大量标注数据,对诈骗话术快速演化和跨场景迁移适应性不足;二是大模型方法多采用端到端判别,缺乏可解释的结构约束,存在幻觉风险;三是事件图谱类研究多偏向风险分析或



模式总结,与具体电话识别任务耦合度有限,难以直接落地应用。上述不足共同制约了诈骗电话识别在小规模样本、新手法和持续演化场景下的性能与可扩展性,急需一种兼具语义理解和结构化表达的新型方法。

针对上述不足,本文提出了一种基于大语言模型与事件演化图谱的诈骗电话识别(fraud call identification based on large language models and event evolution graphs, FLEEG)方法,以在小规模样本和诈骗话术持续演化的条件下,实现更高精度且具备持续扩展能力的诈骗电话识别方法。本文的主要贡献如下。

(1) 基于大语言模型的零样本事件演化链抽取方法:设计专用提示词(Prompt)模板,利用大语言模型的零样本理解与生成能力,直接将主叫侧话术抽象为动宾短语事件链,无需人工标注即可跨案件、跨场景复用。

(2) 事件演化图谱构建与融合方法:以“事件-演化-事件”为骨架,将多条零样本抽取得到的事件链通过语义相似度进行融合、节点合并与向量更新,形成可增量扩展的动态图,从而更清晰地刻画诈骗话术的演化逻辑,并支持对新出现话术模式的持续补充。

(3) 基于事件演化图谱的诈骗电话识别方法:在完全不依赖负样本的情况下,通过事件演化链提取、关键节点度阈值过滤、匹配率计算与连通性检测,提升模型的准确性与可扩展性。

2 FLEEG 构建和诈骗电话识别方法

本文提出了一种基于大语言模型与事件演化图谱的诈骗电话识别方法。该方法借鉴单类分类(one-class learning, OCL)^[20]的思想,采用半监督学习范式,仅使用从历史诈骗通话中提炼的事件链作为正例样本进行建模,从而在无需大量负例数据的情况下实现有效学习。该方法能够判别新通话内容是否符合诈骗事件的演化规律,从而

实现对涉诈电话的识别。

在具体实现过程中,FLEEG方法以“事件-演化-事件”的结构为核心框架。首先,利用大语言模型对每次诈骗通话进行语义解析与抽象,生成对应的事件演化链;随后,基于语义相似度融合相似事件节点,逐步构建高度压缩且可增量扩展的诈骗事件演化图谱。该图谱不仅保留了诈骗行为的时序因果逻辑(如“身份冒充-威胁恐吓-诱导转账”),还通过关键节点之间的连通性检测,实现对风险模式的高效识别。为便于理解,基于大语言模型和事件演化图谱的诈骗电话识别框架如图1所示。图中对事件演化链抽取、诈骗事件演化图谱构建以及诈骗电话识别的处理流程与输出结果进行了展示,并标注了所使用的工具,以帮助读者直观地把握本方法的实现逻辑。

2.1 事件演化链抽取

事件演化链(event evolutionary chain, EEC)用于抽象刻画诈骗通话中主叫方话术的演化过程,其核心思想是将主叫人所采用的话术按时间顺序组织成一组事件序列,可形式化表示为:

$$EEC=[e_1, e_2, e_3, \dots, e_n] \quad (1)$$

其中, e_i 表示第*i*个诈骗事件节点。

具体的构建流程如下:首先,对原始通话内容进行预处理,剔除被叫方的对话信息,仅保留主叫方的话术内容。这样做的目的是在诈骗话术提取过程中消除被叫方的噪声干扰,从而更完整地捕捉诈骗者的引导逻辑与行为链条。在此基础上,构造用于事件抽取的提示词,并调用大语言模型DeepSeek-V3.1-Chat进行自动化抽取,将原始主叫通话内容压缩为不超过10步的事件演化链。例如,[“自称电商客服”,“告知刷单返佣”,“要求下载App”,“说明垫付规则”,“引导高佣单”,“索要税款”,“解释必要性”,“催补缴保证金”,“承诺返款”,“催促结清”]。

这种零样本EEC抽取方法可以有效地将主叫方的意图与话术内容映射为一系列动宾结构的事

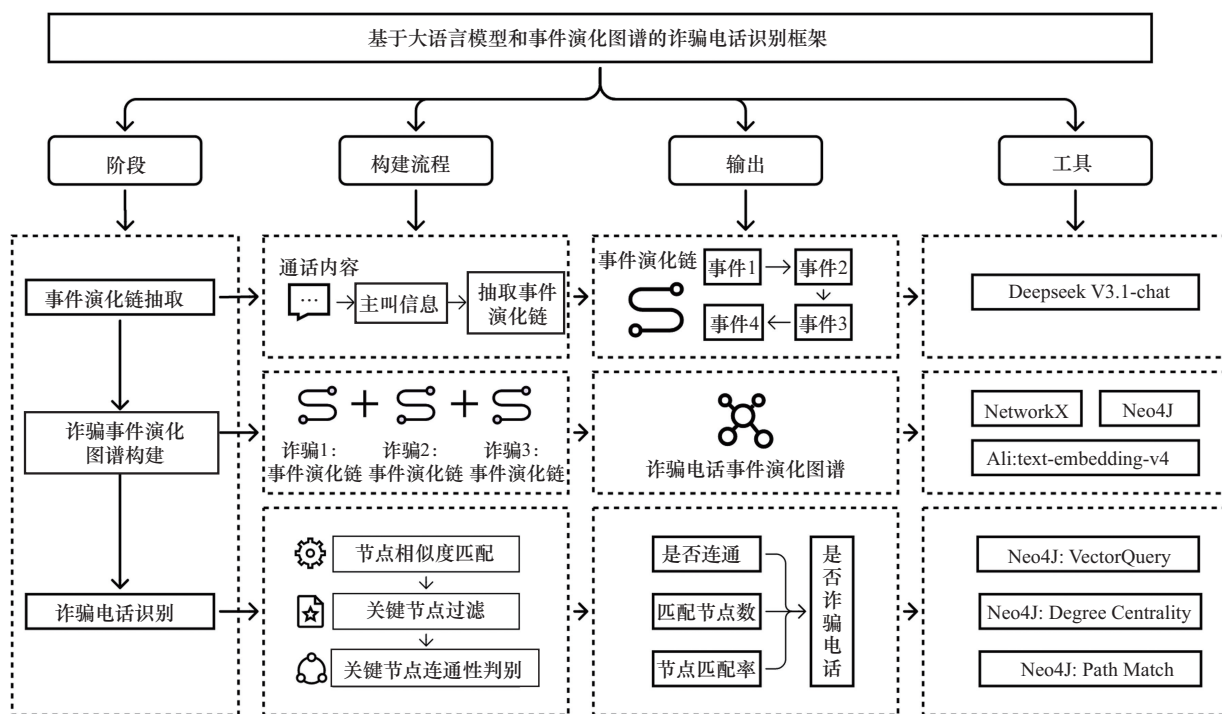


图1 基于大语言模型和事件演化图谱的诈骗电话识别框架

件节点，从而形成一个抽象化、泛化性强且具有时序逻辑的演化过程。与直接基于文本的分析方法相比，EEC不仅更直观地反映了诈骗话术的阶段推进特征，也为后续的事件演化图谱构建和诈骗电话识别提供了标准化输入。EEC抽取提示词模板如图2所示。

【任务】

1. 针对通话主叫人的发言内容，识别并抽取能够体现其意图的核心话术。
2. 将这些话术抽象、泛化为语义完整的动宾短语结构（事件）。
3. 按照通话时间顺序，将事件串联成一个事件演化列表，用来刻画主叫人意。

【要求】

1. 按照通话的时间顺序输出事件，形成事件演化链。
2. 每条链不超过 10 个事件，保留关键转折点即可。
3. 常见事件类型包括：
 - *身份声明（自称身份）
 - *起因说明（告知问题/原因）
 - *诉求提出（要求某种行为）
 - *动机解释（说明理由/强调必要性）
 - *结果承诺（承诺解决问题/带来好处）
 - *行动引导（指引下一步操作）。
4. 严格按照输出格式输出，不要包含无关内容。

【输出格式】

["事件 1(动宾短语)", "事件 2(动宾短语)", "事件 3(动宾短语)"]

【通话主叫信息】

{[content]}

图2 事件演化链抽取提示词模板

2.2 诈骗事件演化图谱构建

诈骗事件演化图谱（fraud event evolution graph, FEEG）可视为一组诈骗事件演化链的压缩与融合结果，是对多条诈骗话术演化路径的统一建模。与单条事件演化链相比，该图谱不仅具备高压缩性的表示能力，而且可随新样本不断更新，具有良好的可扩展性与动态演化特征。

具体而言，首先将前述抽取得到的每一条诈骗电话事件演化链转化为相应的演化图；随后，通过节点语义相似度计算与结构对齐，对这些事件演化图进行逐一合并与融合，从而形成一个整体性的诈骗事件演化图谱。该图谱中的节点以动宾短语表示的事件（event）为基本单元，边则表示事件之间的时序演化关系，因而能够直观刻画诈骗行为的逻辑推进与阶段演变。诈骗事件演化图谱 G 可形式化表示为：

$$G = (V, E, \varphi_v) \quad (2)$$

其中， $V = \{v_1, v_2, v_3, \dots, v_n\}$ 为节点集合， $E \subseteq \{\{u, v\} | u, v \in V, u \neq v\}$ 表示有向边集合， $\varphi_v: V \rightarrow A_v$



为节点属性映射（如事件或向量表示）。

2.2.1 事件演化图构造

在图结构构建方面，本文采用 NetworkX 作为图操作的基础框架，将事件演化链映射为标准化的图数据结构。具体而言，每条事件演化链被转化为一个包含节点与边的有向图。

节点 (node)：以动宾短语形式的事件描述作为节点标识，并以其对应的语义向量作为节点属性 (embedding)，为后续的相似性计算与节点融合提供语义表示支撑。

边 (edge)：按照事件在演化链中的时间顺序依次构建单向边，用以刻画事件之间的时序因果关系。

例如，对于事件演化链：[“自称电商客服”“告知刷单返佣”“要求下载 App”“说明垫付规则”“引导高佣单”“索要税款”“解释必要性”“催补缴保证金”]，其对应的诈骗事件演化图如图 3 所示。该图结构完整保留了诈骗话术从初始身份冒充到逐步诱导转账的演化逻辑，为后续事件演化图谱的融合与全局建模奠定了基础。

2.2.2 事件演化图的合并

演化图的合并是一个递归过程。首先，依据前述方法生成首个事件演化图，作为全局图谱的初始结构。随后，每当抽取到一条新的事件演化链并转化为对应的事件演化图时，便调用图合并算法，将其与当前已构建的大图进行融合。在合并过程中，算法以节点语义相似度为核心判据：

当新图中的节点与大图中已有节点的相似度超过设定阈值时，则判定为同一事件并执行节点对齐；否则，将该节点及其相关边作为新元素增量加入。通过这种方式，图结构能够在保持压缩性的同时不断扩展，逐步形成一个覆盖多类诈骗话术模式的整体性演化图谱。

该递归式合并机制不仅能够在结构层面实现对不同诈骗通话路径的整合，还能通过语义层面的相似性约束保证融合过程的合理性与一致性。其核心的图合并算法如算法 1 所示。

算法 1 图合并算法

输入：两个图 $g_1=(V_1, E_1, \varphi_{v1})$ $g_2=(V_2, E_2, \varphi_{v2})$ ；相似度阈值 τ

输出：合并的大图 G'

(1) 初始图合并

基于 NetworkX 的 Compose 算法： $G' = \text{Compose}(g_1, g_2)$ ，将 g_1 和 g_2 中 ID 相同的节点合并为一个大图 $G'=(V', E', \varphi_{v'})$

(2) 节点相似度计算与合并

对图 g_1 中每个节点 $u \in V_1$ 和图 g_2 中每个节点 $v \in V_2$ ，若 $u \neq v$ 且 $v \in V'$ ，则取节点 u 与 v 的嵌入向量 e_u 与 e_v 计算其余弦相似度 $c_s(u, v) = \cos(e_u, e_v)$ ，其中 c_s 表示两节点间的相似度分值， $\cos(\cdot)$ 表示余弦相似度函数；若 $c_s(u, v) > \tau$ （建议 $\tau = 0.85$ ），则在图 G' 中将节点 u 与 v 合并： $G' = \text{Con}(G', u, v)$ ，并将合并后节点的嵌入向量更新为两者均值： $e_n \leftarrow (e_u + e_v) / 2$ 。

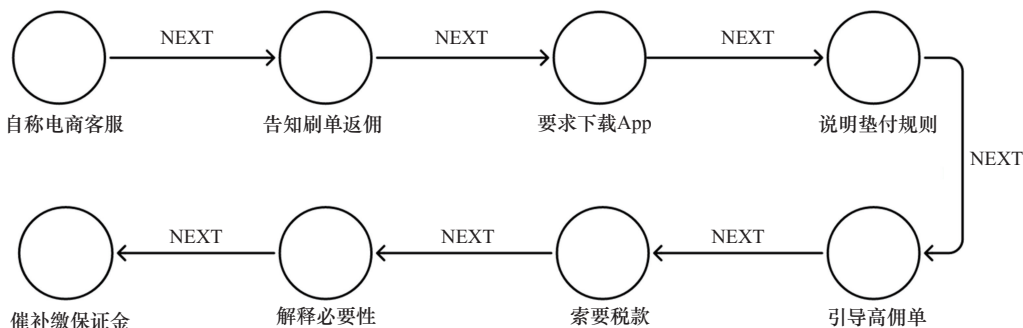


图 3 诈骗事件演化图

(3) 输出合并的大图 G'

将合并后的图 G' 作为下一次合并过程的 g_1 ，不断重复，直到所有图都合并成一个统一的大图 G 。

2.2.3 写入 Neo4j 图数据库

在完成基于 NetworkX 的事件演化图构建后，为实现数据的持久化存储与高效检索，本文将完整的诈骗事件演化图谱写入 Neo4j 图数据库。这样不仅能够长期保存图谱结构，还能充分利用图数据库的索引与查询能力，为后续的图搜索与风险检测提供支撑。

具体流程如下：首先，在 Neo4j 中新建一个欺诈向量索引（`fraud_embedding_idx`），以支持基于余弦相似度的向量检索功能，从而实现对节点向量属性的高效语义检索。随后，对传入的

NetworkX 图 G 进行遍历，将其中的每个节点写入 Neo4j 数据库，并赋予统一的标签 `Fraud`，同时存储节点标识（`node id`）及其对应的向量作为节点属性。最后，遍历图中的边集合，在 Neo4j 数据库中为相关节点建立 NEXT 类型的有向关系，用以表示事件之间的时序演化依赖。至此，整个图结构便完成了从内存中的 NetworkX 表示到数据库中 Neo4j 图谱的写入与重建。在写入 Neo4j 数据库后，通过 Explore 工具搜索得到的图谱局部结构示例如图 4 所示。

至此，诈骗事件演化图谱的构建过程完成。值得注意的是，该方法仅依赖诈骗电话类样本即可实现建模，无须引入正常通话数据，从而有效避免了负样本收集困难的问题。基于这一图谱，能够获得对诈骗话术演化模式的结构化表示。下一小节将进一步阐述如何利用所构建的诈骗事件

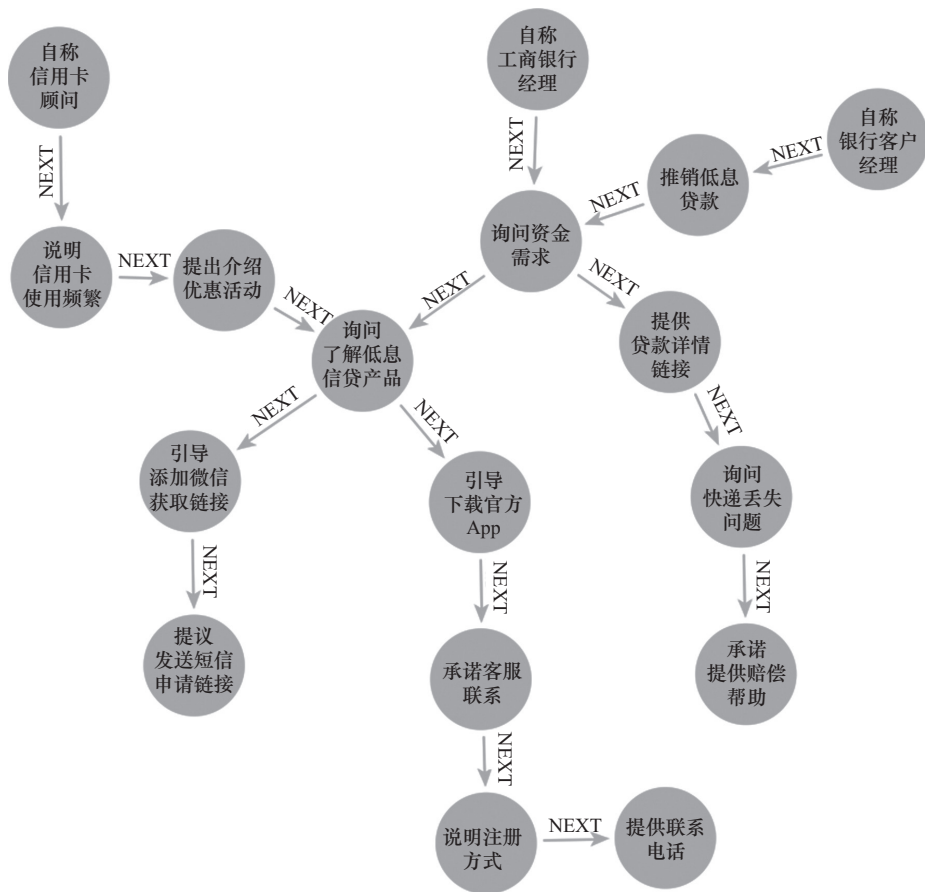


图4 诈骗事件演化图谱局部结构示例



演化图谱开展诈骗电话识别任务。

2.3 诈骗电话识别

诈骗电话的识别过程包括4个主要步骤，具体如下。

步骤1 对输入通话的主叫内容进行处理，利用大语言模型抽取其对应的EEC。

步骤2 将该EEC与已构建的诈骗事件演化图谱进行节点相似度匹配，以识别通话事件与图谱中已知诈骗模式的关联性。

步骤3 对匹配结果进行筛选，剔除语义贡献较弱或处于边缘位置的非关键节点，仅保留能够反映诈骗演化逻辑的核心节点。

步骤4 对剩余关键节点进行连通性判别，并结合关键节点数量、关键节点匹配率以及连通性指标进行综合评估，最终输出该通话是否属于诈骗电话的判定结果。

2.3.1 事件演化链提取

首先，输入待处理的通话文本，并对其进行预处理，去除其中的被叫方信息，仅保留主叫方的话术内容，以确保后续抽取过程专注于诈骗者的行为逻辑；然后，基于第2.1节所提出的大语言模型事件演化链抽取方法，对主叫方通话内容进行解析，生成对应的EEC；随后，将该EEC进一步转化为事件演化图 g ，以获得结构化表示。

2.3.2 节点相似度匹配

设待匹配的图 g 的节点为 $\text{Nodes}=\{n_1, n_2, n_3, \dots, n_k\}$ ，对应节点向量集合为 $\mathbf{Q}=\{q_1, q_2, q_3, \dots, q_k\}$ ，诈骗事件演化图谱 G 向量索引中的节点集合为 $\mathbf{V}=\{v_1, v_2, v_3, \dots, v_n\}$ ，相似度函数为 $\text{sim}(q, v)$ （如余弦相似度），相似度阈值为 τ 。

首先，在向量索引中对每个查询向量 $q \in \mathbf{Q}$ 寻找Top K （建议 $K=1$ ）最相似的节点 $v \in \mathbf{V}$ ，并记录其相似度 s ，得到匹配节点集合 $\mathbf{M}$ ：

$$\mathbf{M} = \bigcup_{q \in \mathbf{Q}} \{v, s \mid v \in \text{Top}K(\text{sim}(q, v)), s = \text{sim}(q, v)\} \quad (3)$$

随后，基于相似度阈值 τ （建议 $\tau=0.91$ ）对匹配节点进行筛选，只保留相似度高于阈值的节点：

$$\mathbf{M}_n = \{v \mid (v, s) \in \mathbf{M} \wedge s \geq \tau\} \quad (4)$$

最终输出的 $\mathbf{M}_n$ 即为满足条件的节点集合，作为后续关键节点筛选与连通性判别的基础。

2.3.3 关键节点过滤

对于已匹配的节点，需要进行关键节点筛选。在此过程中，将根据节点的重要性指标保留对诈骗事件判定具有关键作用的节点，同时忽略语义贡献较低或冗余的非关键节点。基于筛选后的关键节点，可进一步计算其在候选节点集合中的匹配率，以量化关键事件的覆盖程度。具体操作流程如下。

采用度中心性（degree centrality）作为关键节点识别方法。设节点的入度为 $d^{\text{in}}(v)$ ，出度为 $d^{\text{out}}(v)$ ，则节点的总度数定义为：

$$C_D(v) = d^{\text{in}}(v) + d^{\text{out}}(v) \quad (5)$$

节点总度数通常用于衡量其在网络中的连接程度，在图分析中常用于识别高连接节点，例如，高风险账户、热门商品或关键中介节点等。

在本文方法中，对序列中的每个匹配节点依次计算其入度与出度之和，将结果加入度数列表，从而得到整个序列的节点度分布。随后，设定阈值 d ，将总度数 $\geq d$ 的节点作为关键节点 K_n 输出，即

$$K_n = \{v \in \mathbf{M}_n \mid C_D(v) \geq d\} \quad (6)$$

通过该策略，可有效筛选出在诈骗事件演化图中结构关键且对风险判定具有重要意义的核心节点，为后续的连通性判别和综合风险评估提供依据。

2.3.4 节点连通性判别

在获得过滤后的关键节点集合后，需要在事件演化图谱 G 中进行连通性判别，以验证这些关

键节点是否能够通过有限跳数形成连贯的演化路径。具体方法如下。

首先，输入关键节点序列 $K_n=(v_1, v_2,$

$v_3, \dots, v_k)$ 、最大跳数参数 h （建议取 $h=4$ ）。

随后，判定边路径存在性。对于序列中相邻的节点对 (v_i, v_{i+1}) ，定义布尔变量 q_i 如下：

$$q_i = \begin{cases} 1, & \text{在 } G \text{ 中存在从 } v_i \text{ 到 } v_{i+1} \text{ 的路径且路径长度 } \leq h \\ 0, & \text{其他} \end{cases} \quad \forall_i \in \{1, 2, \dots, k-1\} \quad (7)$$

即，当且仅当在图谱 G 中存在从 v_i 到 v_{i+1} 的路径，且路径长度不超 h 时， $q_i=1$ ；否则， q_i 为 0。

接着，判定整体连通性。通过对所有相邻节点对的连通性进行逻辑与操作，得到整体结果：

$$R = \bigwedge_{i=1}^{k-1} q_i \quad (8)$$

最后，输出判定结果。若 $R=1$ ，则认为关键节点序列在图谱中连通；否则，判定为不连通。关键节点连通性 C 可表示为：

$$C = \begin{cases} \text{True}, & R=1 \\ \text{False}, & R \neq 1 \end{cases} \quad (9)$$

该方法能够有效评估关键节点在诈骗事件演化图谱中的时序可达性，为后续的诈骗电话判别提供结构化依据。

2.3.5 诈骗电话判别

在完成关键节点筛选与连通性判别后，进一步引入以下 3 个指标对通话进行综合评估。

(1) 关键节点数： $|K_n|$ ，即筛选后关键节点的数量。

(2) 关键节点匹配率： $M_r = \frac{|K_n|}{|M_n|}$ 。其中， $|M_n|$ 表示候选匹配节点总数。

(3) 关键节点连通性： C ，表示关键节点是否在图谱中满足可达性约束。

基于上述指标，最终的诈骗电话识别规则可定义如下（ m 建议为 3， r 建议为 0.25）：

$$\text{Output} = \begin{cases} \text{True}, & C(K_n) \wedge |K_n| \geq m \wedge M_r > r \\ \text{False}, & \text{其他} \end{cases} \quad (10)$$

3 实验与结果分析

3.1 实验数据

3.1.1 数据来源

本文实验所使用的数据主要来源于两个渠道：（1）TeleAntiFraud-28k 数据集^[13]。该数据集基于真实诈骗电话场景，通过大语言模型与多智能体协作方式生成，共包含 16 819 条数据。需要说明的是，该数据集并未公开原始的诈骗电话文本内容，仅提供了基于其加工后的数据。为补充原始诈骗电话数据规模，本文引入下述数据集。（2）CCL23-Eval 任务 6 总结报告——电信网络诈骗案件分类（CCL23-Eval-Task6）^[21]。该任务开源了来自公安部门反诈大数据平台的 102 762 条电诈案情数据。本文基于这些真实案件描述进一步生成了相应的诈骗电话语料，作为实验的重要补充。

实验所使用的实际数据整体分布情况见表 1。TeleAntiFraud-28k 数据集共包含 16 819 条样本，其中包含电话诈骗信息 7 581 条，非电话诈骗信息 9 238 条。CCL23-Eval-Task6 数据集共收录 102 762 条诈骗案件信息，由于其中包含了非电话类诈骗案件，本文采用启发式的正则表达式规则对数据进行初步过滤，最终获得 15 312 条电话诈骗案件信息。

表 1 实验所使用的实际数据整体分布情况 单位：条

数据集	电话诈骗案件信息	非电话诈骗案件信息
TeleAntiFraud-28k	7 581	9 238
CCL23-Eval-Task6	15 312	0

**【角色】**

你是一名反电话诈骗专家。

【任务】

根据提供的案情信息，模拟诈骗者与受害者之间的首次电话通话内容，用于识别类似诈骗电话。

【要求】

1. 诈骗者用 A 表示，受害者用 B 表示。
2. 如果案情信息反映的不是电话诈骗，直接返回“非电话诈骗”。
3. 真实还原当时的通话内容。
4. 输出内容仅包含模拟的通话，不做任何解释或额外说明。

【案情信息】

{{content}}

图5 诈骗电话生成提示词

基于上述两类数据来源，本文在保证数据真实性和多样性的同时，兼顾了场景覆盖范围与实验可复现性，为后续模型训练与评估奠定了坚实基础。

3.1.2 诈骗电话生成

在 CCL23-Eval-Task6 数据集中，随机抽取了 5 000 条案件描述数据，并利用大语言模型生成对应的诈骗者首次通话内容。最终共生成 4 890 条通话数据，其中有 110 条因被模型判定为“非电话诈骗”案件而被舍弃。用于将报案信息转化为首呼诈骗电话的提示词设计如图 5 所示。

3.1.3 数据集划分

基于上述通话数据，本文将其划分为训练集与评估集。训练集共包含 8 000 条电话诈骗样本；评估集共 3 715 条样本，包括 1 890 条电话诈骗样本与 1 825 条非电话诈骗样本。整体来看，训练集占比为 68.29%，评估集占比为 31.71%。各类别数据的详细分布情况见表 2。

表2 训练集和评估基数据的详细分布情况 单位：条

数据集	训练集	评估集 (涉诈)	评估集 (非涉诈)	合计
TeleAntiFraud-28k	4 054	946	1 825	6 825
CCL23-Eval-Task6	3 946	944	0	4 890
合计	8 000	1 890	1 825	11 715

3.2 实验设置

本文提出的方法融合了大语言模型与事件演化图谱技术，并采用近似 OCL 的建模范式。在实验环节中，主要设计了两类对比方案：一类是直接基于大语言模型的识别方法，另一类是结合 OCL 思想的传统机器学习方法。通过对比两类方法在相同数据集上的表现，可以全面验证本文所提方法的有效性与相对优势。

实验评估主要从以下两个方面进行：（1）缩放律（Scaling Law）。记录模型在不同图参数规模下的性能指标，分析模型性能随参数量变化的规律性，以验证实验方法在大规模模型下的可扩展性。（2）准确性。参考传统分类指标，包括精准率（Precision）、召回率（Recall）以及 $F1$ 值，用于量化模型在标准测试集上的整体识别效果。

通过上述多维度评估，可以全面刻画本文方法在实际诈骗电话识别任务中的性能表现及可扩展性，为方法的有效性和适用性提供充分证据。相关代码与数据已开源，在 github 官网搜索 FLEEG 可见。

3.2.1 FLEEG 模型设置

在 FLEEG 模型实验中，训练样本规模由 1 000 条逐步扩展至 8 000 条，依次生成 FLEEG-

1k至FLEEG-8k共8个版本的图谱。上述实验设计可以系统评估图谱规模对诈骗电话识别性能的影响,并进一步分析模型在不同训练样本量下的缩放规律。随后,结合第2.3节所述的诈骗电话识别方法,对各图谱进行系统测试与性能评估。实验参数设置如下。

图谱构建参数:事件描述向量维度设为512,节点合并相似度阈值为0.85。

诈骗电话识别参数:节点匹配相似度阈值为0.91,关键节点度数阈值为5,关键节点匹配数阈值为3,关键节点匹配率阈值为0.25,连通性检测最大跳数为4。

3.2.2 大模型设置

文献[13]针对大模型在诈骗电话识别任务上的表现进行了系统测试,并给出了详尽的实验结论。在对比实验中直接采用文献[13]的评测结果,对多种大模型在诈骗检测任务上的表现进行比较。对比范围涵盖:(1)大语言模型,包括Qwen2.5-72B、DeepSeek-V3、Doubao-1.5-Pro、DeepSeek-R1。(2)多模态大模型,包括GPT-4o、Gemini-2-Flash、Qwen2-Audio-7B-Instruct。(3)专项模型,即文献[13]中基于Qwen2-Audio-7B-Instruct微调得到的AntiFraud-Qwen2-Audio多模态模型。

3.2.3 OCL模型设置

本文方法借鉴了OCL的思想,选择了两类经典的OCL算法作为对比基线模型:One-Class SVM与Isolation Forest。这两种算法广泛应用于垃圾邮件过滤、欺诈检测等异常检测场景,而诈骗电话识别同样可被视为异常检测问题,因此二者在本研究中具有较高的可比性。

具体建模流程如下。

(1)使用阿里巴巴提供的text-embedding-v4模型,将电话主叫信息统一转换为512维向量表示。

(2)基于这些向量分别训练One-Class SVM与Isolation Forest模型,并将训练样本规模从

1 000条逐步扩展至8 000条。

(3)对两个模型的关键超参数 μ (One-Class SVM)与contamination(Isolation Forest)进行网格搜索,参数范围设置为[0.2,0.1,0.05]。

(4)以F1值最高的超参数配置作为最终实验结果。

通过上述设置,可系统评估传统OCL算法在诈骗电话识别任务中的性能表现,并与本文提出的方法进行对比分析。

3.3 实验结果

3.3.1 模型缩放律

FLEEG模型缩放律如图6所示。随着训练样本规模的扩展,FLEEG模型的图谱节点规模(S Node)由2 965增长至12 336。与此同时,模型的精准率(P)由98.55%小幅下降至93.48%,但召回率(R)与F1值均呈持续上升趋势:召回率由61.27%提升至89.52%,F1值则由75.56%增长至91.46%。

作为对比,One-Class SVM与Isolation Forest两类传统OCL模型在训练样本规模增加的情况下,其F1值并未表现出明显提升,整体保持相对稳定。

综上结果表明,FLEEG模型能够随着图规模的增加不断提升性能,符合缩放律的特征。

3.3.2 模型性能

模型性能指标对比见表3。从表3的对比评测结果来看,不同类型模型在诈骗电话识别任务中的表现差异较为明显。大语言模型整体表现中规中矩,其中表现最佳的是具备思维链机制(Chain-of-Thought, CoT)的DeepSeek-R1,其F1值达79.25%。多模态大模型在未经过微调时效果较弱,而经过任务定向微调后的AntiFraud-Qwen2-Audio性能显著提升,F1值提升至84.78%。在OCL基线模型中,尽管One-Class SVM在5 000样本规模下取得了73.17%的F1值,Isolation Forest在7 000样本规模时达到71.86%,

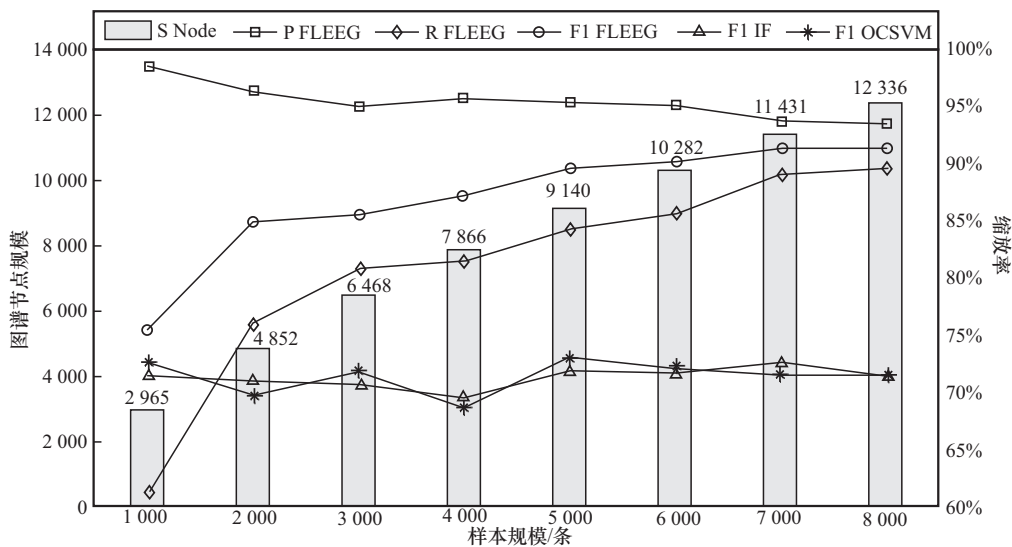


图6 FLEEG模型缩放律

表3 模型性能指标对比

模型类型	模型名称	精准率	召回率	F1 值
大语言模型	Qwen2.5-72B	—	—	51.44%
	Doubao-1.5-Pro	—	—	36.11%
	DeepSeek-V3	—	—	14.62%
	DeepSeek-R1	—	—	79.25%
多模态大模型	GPT-4o	—	—	50.00%
	Gemini-2-Flash	—	—	59.61%
专项模型	Qwen2-Audio-7B-Instruct	—	—	58.51%
	AntiFraud-Qwen2-Audio	—	—	84.78%
OCL 基线模型	One-Class SVM-5 000	59.10%	96.03%	73.17%
	IsolationForest-7 000	60.82%	90.11%	72.62%
FLEEG	FLEEG-1 000	98.55%	61.27%	75.56%
	FLEEG-2 000	96.31%	75.98%	84.94%
	FLEEG-3 000	95.08%	80.74%	85.6%
	FLEEG-4 000	95.77%	81.53%	87.22%
	FLEEG-5 000	95.39%	84.34%	89.52%
	FLEEG-6 000	95.12%	85.66%	90.14%
	FLEEG-7 000	93.71%	89.1%	91.35%
	FLEEG-8 000	93.48%	89.52%	91.46%

但整体精度偏低，误报率较高。

相比之下，FLEEG 模型展现出显著优势。在仅使用 2 000 小规模样本时， $F1$ 值便达到 84.94%，已与当前表现最优的 AntiFraud-Qwen2-Audio 相当。这一结果表明，FLEEG 方法不仅具备较强的泛化能力，能够有效归纳和总结不同诈

骗话术的共性模式，而且在小规模样本条件下也能保持稳定的性能表现。更为重要的是，随着训练样本规模的增加，FLEEG 的性能持续提升，充分体现了良好的缩放律特性。在 8 000 样本规模下，FLEEG 的 $F1$ 值进一步攀升至 91.46%，显著优于对比模型。

4 结束语

本文针对电信诈骗话术持续演化、标注成本高等问题,提出了一种融合大语言模型语义能力与事件演化图谱结构约束的诈骗电话识别思路。该方法从话术演化过程出发,将诈骗行为建模为可扩展的事件演化结构,在无需负样本的条件下实现对诈骗话术的有效识别。实验结果表明,该方法在小样本和规模扩展场景下均表现出良好的稳定性与可扩展性,事件级结构化建模能够在一定程度上缓解大语言模型的幻觉风险问题,为持续演化诈骗话术的结构化建模与识别提供了一条可行的技术路径。

尽管本文的研究取得了一定效果,但当前方法在召回率方面仍存在提升空间,在强调最大化可疑诈骗发现的应用场景下仍具有一定局限。未来研究可通过进一步优化判别条件(如关键节点匹配数量与匹配率阈值),引入“疑似诈骗”中间类别与人工审核机制,以提升整体召回能力和实际适用性。

总体而言,本研究表明,将事件演化图谱引入大语言模型驱动识别框架,有助于在复杂、持续演化的行为识别任务中实现语义理解与结构约束的协同建模,为相关问题的研究提供了新的思路和方法参考。

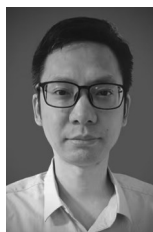
参考文献:

- [1] 陈禹潼. 刑事一体化视野中的跨境电信网络诈骗犯罪疑难问题研究[J]. 法律适用, 2024(10): 117-132.
Chen Y T. Study on the difficulties of regulation on cross-border telecom and online fraud crimes in the perspective of criminal integration[J]. Journal of Law Application, 2024(10): 117-132.
- [2] 单勇. 《反电信网络诈骗法》的均衡实现: 预防性秩序和正当性权益之间[J]. 法制与社会发展, 2025, 31(1): 40-64.
Shan Y. Balanced realization of anti-telecom and online fraud law: between preventive order and legitimate rights[J]. Law and Social Development, 2025, 31(1): 40-64.
- [3] 张照星, 胡琨璇, 范英, 等. 电信诈骗案件中基于运营商国际漫游话单的嫌疑号码拓展[J]. 信息安全, 2021, 21(S1): 12-16.
Zhang Z X, Hu C X, Fan Y, et al. The expansion of suspect number based on operator's international roaming bill in telecom fraud cases[J]. Netinfo Security, 2021, 21(S1): 12-16.
- [4] 李洋, 李振华, 辛显龙. 基于攻击经济学的移动虚拟运营商诈骗检测[J]. 计算机科学, 2023, 50(8): 260-270.
Li Y, Li Z H, Xin X L. Attack economics based fraud detection for MVNO[J]. Computer Science, 2023, 50(8): 260-270.
- [5] 刘晨迪, 李耀, 陈佳逸, 等. 融合图结构和多通道注意力机制的诈骗电话识别[J]. 计算机应用, 2025, 45(S2): 58-63.
Liu C D, Li Y, Chen J Y, et al. Scam call identification by integrating graph structure and multi-channel attention mechanism[J]. Journal of Computer Applications, 2025, 45(S2): 58-63.
- [6] 周俊杰, 许鸿奎, 卢江坤, 等. 引入位置信息和 Attention 机制的诈骗电话文本分类[J]. 小型微型计算机系统, 2023, 44(11): 2502-2509.
Zhou J J, Xu H K, Lu J K, et al. Position embedding and attention are introduced into the fraudulent phone text classification[J]. Journal of Chinese Computer Systems, 2023, 44(11): 2502-2509.
- [7] 司海平, 李阔, 李婷婷, 等. 基于提示对比学习的小样本电信诈骗文本分类方法研究[J]. 计算机应用与软件, 2025: 1-8. (2025-01-27).
Si H P, Li K, Li T T, et al. Research on few-shot telecom fraud text classification approaches based on prompt-driven contrastive learning[J]. Computer Applications and Software, 2025: 1-8. (2025-01-27).
- [8] 邓诗琦, 洪亮. 面向智能应用的领域本体构建研究: 以反电信诈骗领域为例[J]. 数据分析与知识发现, 2019, 3(7): 73-84.
Deng S Q, Hong L. Constructing domain ontology for intelligent applications: case study of anti tele-fraud[J]. Data Analysis and Knowledge Discovery, 2019, 3(7): 73-84.
- [9] 刘卓娴, 石拓, 胡啸峰. 电信网络诈骗警情要素抽取与模式挖掘研究[J]. 情报杂志, 2025, 44(6): 168-176, 192.
Liu Z X, Shi T, Hu X F. Research on alarm texts element extraction and pattern mining of telecom network fraud[J]. Journal of Intelligence, 2025, 44(6): 168-176, 192.
- [10] 李衡峰, 郑榕, 邓菁, 等. 基于大模型的电信网络诈骗预警技术研究[J]. 警察技术, 2025(4): 13-18.
Li H F, Zheng R, Deng J, et al. Research on telecom network fraud early warning technology based on large language model[J]. Police Technology, 2025(4): 13-18.
- [11] Shen Z T, Wang K Z, Zhang Y Q, et al. Combating phone scams with LLM-based detection: where do we stand? (student



- abstract)[J]. Proceedings of the AAAI Conference on Artificial Intelligence, 2025, 39(28): 29487-29489.
- [12] Singh G, Singh P, Singh M. Advanced real-time fraud detection using RAG-based LLMs[PP]. V1. arXiv (2025-01-25) [2025-09-23]. arXiv: 2501.15290.
- [13] Ma Z M, Wang P D, Huang M H, et al. TeleAntiFraud-28k: an audio-text slow-thinking dataset for telecom fraud detection[PP]. V4. arXiv (2025-08-18)[2025-09-23]. arXiv: 2503.24115.
- [14] 裴炳森, 李欣, 吴越. 基于 ChatGPT 的电信诈骗案件类型影响力评估[J]. 计算机科学与探索, 2023, 17(10): 2413-2425.
- Pei B S, Li X, Wu Y. Influence evaluation of telecom fraud case types based on ChatGPT[J]. Journal of Frontiers of Computer Science and Technology, 2023, 17(10): 2413-2425.
- [15] Guan S P, Cheng X Q, Bai L, et al. What is event knowledge graph: a survey[J]. IEEE Transactions on Knowledge and Data Engineering, 2023, 35(7): 7569-7589.
- [16] Li Z Y, Zhao S D, Ding X, et al. EEG: knowledge base for event evolutionary principles and patterns[M]//Social Media Processing: Communications in Computer and Information Science. Singapore: Springer Singapore, 2017: 40-52.
- [17] Ding X, Li Z Y, Liu T, et al. ELG: an event logic graph[PP]. V2. arXiv (2019-08-07)[2025-09-23]. arXiv: 1907.08015.
- [18] 周胜利, 徐睿, 陈庭贵, 等. 生成式人工智能驱动下电信网络诈骗风险演化实证研究[J]. 电信科学, 2025, 41(5): 149-165.
- Zhou S L, Xu R, Chen T G, et al. Empirical study on the evolution of telecom fraud risks driven by artificial intelligence generated content [J]. Telecommunications Science, 2025, 41(5): 149-165.
- [19] 斯彬洲, 孙海春, 吴越. 基于大语言模型和事件融合的电信诈骗事件风险分析[J]. 数据分析与知识发现, 2025, 9(7): 38-51.
- Si B Z, Sun H C, Wu Y. Risk analysis of telecom fraud events based on large language models and event fusion[J]. Data Analysis and Knowledge Discovery, 2025, 9(7): 38-51.
- [20] Perera P, Oza P, Patel V M. One-class classification: a survey[PP]. V1. arXiv (2021-01-08)[2025-09-23]. arXiv: 2101.03064.
- [21] Sun C J, Ji J, Shang B Y, et al. Overview of CCL23-Eval Task 6: telecom network fraud case classification[C]//Proceedings of the 22nd Chinese National Conference on Computational Linguistics (Volume 3: Evaluations). Harbin: Chinese Information Processing Society of China, 2023: 193-200.

[作者简介]



胡程忆 (1980-), 男, 现就职于卓望信息技术 (北京) 有限公司, 主要从事自然语言处理与安全风控工作。