

基于知识图谱增强的大语言模型预检分诊框架研究

刘会 姚宽达 刘湘闽 胡佳慧 王茜 方安

(中国医学科学院北京协和医学院, 医学信息研究所 北京 100020)

[摘要] 目的/意义 基于知识图谱增强的大语言模型技术构建预检分诊框架, 以提升预检分诊的智能化水平, 改善患者就诊体验。**方法/过程** 结合知识图谱与医院科室目录信息, 提出一种基于知识图谱增强大语言模型预检分诊框架 (KGE-LLM TF), 将患者症状描述精准映射至目标医院的具体科室, 实现通用知识图谱与异构医院挂号系统间的自适应对接。**结果/结论** 实验表明, 该框架方法的 **Top-1 准确率 (ACC@1)** 与 **Top-3 准确率 (ACC@3)** 分别达到 **84.91%**与 **92.45%**, 复杂案例 **F1** 值优于其他基线模型 **25%**以上, 能够有效处理症状描述模糊、科室名称异构等现实挑战, 为智能化预检分诊提供可靠支撑。**[关键词]** 医学知识图谱; 知识增强; 大语言模型; 智能分诊

Research on pre-examination and triage framework of large language model based on knowledge graph enhancement

LIU Hui, YAO Kuanda, LIU Xiangmin, HU Jiahui, WANG Qian, FANG An
Institute of Medical Information, Chinese Academy of Medical Sciences & Peking Union Medical College, Beijing 100020, China

[Abstract] Purpose/Significance The pre-examination and triage framework is constructed based on the large language model technology enhanced by knowledge graph to improve the intelligence level of pre-examination and triage and improve the patient 's experience. **Method/Process** Combining knowledge graph and hospital department directory information, a knowledge graph enhanced large language model pre-examination and triage framework (KGE-LLM TF) is proposed to accurately map the patient 's symptom description to the specific departments of the target hospital, so as to realize the adaptive docking between the general knowledge graph and the heterogeneous hospital registration system. **Result/Conclusion** Experiments show that the **Top-1 accuracy (ACC@1)** and **Top-3 accuracy (ACC@3)** of the framework method reach 84.91 % and 92.45 % respectively, and the F1 value of complex cases is better than that of other baseline models by more than 25 %. It can effectively deal with practical challenges such as vague symptom description and heterogeneous department names, and provide reliable support for intelligent pre-examination and triage.

[Keywords] Medical Knowledge Graph; Knowledge Enhancement; Large Language Model; Intelligent Triage

1 引言

预检分诊是优化资源配置、保障医疗安全的关键环节^[1]。尽管患者对健康服务的精准性与时效性提出了更高要求^[2], 但在实际就医场景中, 医学术语的专业壁垒、患者自我认知的模糊性、综合性医院门诊量增长与专科细化等, 常导致挂号误选, 不仅降低就诊效率更造成医疗资源的错配与浪费。虽然基于开放数据的通用知识图谱已初步实现了“症状-疾

[修回日期] 2026-4-30

[作者简介] 刘会, 中级工程师; 通信作者: 方安, 研究馆员。

[基金项目] 中国医学科学院医学与健康科技创新工程项目“医学知识管理与智能化知识服务关键技术研究”; 中央高水平医院临床科研专项“以 ICD-11 和欧盟 Orphanet 罕见病分类体系为基础的中国罕见病分类标准体系构建研究”

病-科室”的关联建模，但在跨机构的实际部署中，通用知识与异构业务系统之间存在显著的“语义鸿沟”：不同医疗机构在科室命名（如“消化内科”与“胃肠科”）、亚专科划分及服务范围上的差异，使得通用模型难以直接适配具体医院的挂号目录。

针对上述工程落地难题，本研究设计并实现了一个知识图谱增强的大语言模型协同框架，构建了一种任务驱动的动态对齐机制，能够在不依赖特定医院信息系统（HIS）改造的前提下，分析推理患者病症并适配目标医院的科室体系。本研究提供了一种高鲁棒性、易迁移的患者自助分诊解决方案，强调工程可行性与部署灵活性，推动了智能分诊从“通用知识库查询”向“场景化业务适配”的范式转变，为实现低成本、广覆盖的智慧医疗服务提供了可行的技术路径。

2 相关研究现状

2.1 预检分诊

当前研究表明，智能化系统在优化预检分诊流程、提升关键绩效指标方面已取得实证成效。例如，刘徐方等^[3]验证了智能化系统在急诊场景中的应用，其在分诊准确率、抢救成功率及患者满意度方面均显著优于常规方法，同时有效缩短了候诊时间并降低了医疗风险。苗开贵等^[4]的研究进一步证实，智能系统管理有助于提升门诊分诊效率与安全性。上述研究多采用规则引擎或统计模型，优化目标集中于单一医院内部的接诊流程，智能化赋能主要面向分诊护士，而非直接服务患者。患者就诊挂号时，仍需自行解读不同医院的科室命名与挂号规则，流程效率增益被显著削弱。

2.2 医学知识图谱

知识图谱是由具有属性的实体通过关系连接而成的网状知识库^[5]。依托海量医学数据资源及所构建的医学知识图谱，可以为患者、医疗服务者等提供医疗辅助决策支持^[6]。医学知识图谱作为领域特定的结构化知识表示，其知识关系复杂（存在1-1、1-N、N-1、N-N等多种映射）^[7]、语义歧义性强且容错率要求极高。近年来，医学领域涌现了一系列重要成果。其中，中国医学科学院医学信息研究所构建的 MedKaaS^[14] (Medical Knowledge as a Service) 提供了标准化、自动化的构建流程与工具，为实现高质量、可复用的医学知识服务提供了方法论支持^[15]。本研究基于 MedKaaS 系统构建以“症状-疾病-科室”为核心的任务导向轻量化子图谱，显式强化“症状-疾病-通用科室”三元组，并将通用科室作为桥接节点，使其能够与任意目标医院的科室目录进行动态对齐。这种设计为后续的实体映射提供了结构化的“中间语义层”，降低了从通用知识到具体部署的推理难度。

2.3 大语言模型在医学领域中的应用

近年来，大语言模型（LLMs）在医疗诊断推理与决策支持中潜力显著，但受限于“幻觉”问题与知识滞后性。检索增强生成（RAG, Retrieval-Augmented Generation）通过动态注入外部权威知识，成为构建可靠医疗人工智能系统的关键^[16]。临床决策中已有研究者将其应用于放射科流程简化^[17]、慢性肾病指南咨询等场景，减少模型幻觉及不相干内容的输出^[18]。并且，Cheng Ye 等人整合知识图谱嵌入模型，利用知识图谱增强 RAG 系统^[19]，显著提高了药物发现的预测准确性^[20]。Feng Y 等人将大模型与知识图谱相结合，减少推理中的事实性错误，通过潜在的药物-癌症关联促进现有药物新用途的发现^[21]。此外，一个基于知识图谱增强技术构建的高血压性脑出血患者诊断治疗决策支持系统的总体准确率为 94.87%，有潜力为神经外科医生提供快速可靠的决策支持^[22]。王博等^[23]介绍了通过有监督医疗指令微调构建症状-诊断、诊断-治疗等多任务医疗大语言模型方法。

尽管上述 KG 增强 LLM 范式在单病种诊断与用药推荐中成效显著，但在预检分诊场景下仍存在明显局限：第一，多聚焦单一疾病领域（如高血压脑出血、药物靶点发现），缺乏在多科室、多病种综合性任务上验证；第二，知识图谱与术语体系通常同源，未解决通用知识与异构医院科室列表的语义鸿沟。为此，本研究将知识图谱增强 LLM 范式应用于患者自助式跨机构预检分诊，重点探索在真实科室命名体系下的自适应推理，以实现通用知识向具体院务服务的精准适配。

2.4 语义匹配与实体对齐技术

在通用医学术语与异构医院科室名称的精准映射中，核心挑战在于实体对齐与语义匹配。早期方法依赖字符串匹配或词典规则，效率高但泛化性差，难以处理同义异形表述。随后，知识图谱嵌入（如 TransE^[24]、RotatE^[25]）通过低维向量计算语义相似度，虽能捕捉部分语义关联，但受限于图谱完备性及对文本细微差异的敏感度不足。以 BERT^[26]为代表的预训练语言模型有强上下文感知能力，能够理解“胸痛门诊”与“心内科”之间的深层语义关联，显著提升了复杂语义匹配精度^[27]，Sentence-BERT^[28]等进一步优化了句子级语义相似度计算，成为主流技术基础。

尽管上述方法在规范数据上表现良好，但多局限于静态通用语义匹配，难以支撑需结合疾病知识、科室职责与患者特征的临床推理（如“胸痛伴大汗”应匹配“心内科”、“急诊科”或“胸痛中心”）。对此，本研究提出任务驱动的动态对齐框架，其核心区别

在于：（1）根据患者症状描述动态调整对齐策略，替代预计算固定相似度；（2）将目标医院科室列表作为输出约束，而非直接匹配目标；（3）借助知识图谱中介知识（如“冠心病→心血管内科”）实现跨术语体系桥接。通过融合患者个体情况、通用医学知识与目标医院具体目录的动态上下文，构建任务感知的推理框架，根据每次输入的患者症状和科室目录动态生成推理路径，显著提升真实场景下的准确率与适应性。

3 研究方法

3.1 知识图谱构建方法

基于 MedKaaS 提供的标准化、自动化的流程和工具构建图谱，通过数据导入、本体配置、知识抽取及知识融合等步骤，从非结构化文本中抽取相关的知识实体及实体关系，并通过本体定义各类实体之间的关系。以实体为节点、关系为链接，可以形成<实体,关系,实体>的三元组，利用 Nebula 图数据库对三元组进行存储，最终形成知识图谱。

3.2 基于知识增强的大语言模型预检分诊框架

为解决通用知识图谱与具体医院挂号场景间的“语义鸿沟”问题，实现精准的挂号科室推荐，本节提出一种知识增强的大语言模型预检分诊框架（KGE-LLM TF, Knowledge Graph-Enhanced Large Language Model Triage Framework）。其核心思想是：利用通用知识图谱提供结构化、可信的医学先验知识，引导大语言模型理解患者症状的临床含义，并将其与目标医院特有的、异构的科室目录进行上下文关联与推理，最终输出符合该医院具体语境的最可能科室。通过显式引入结构化知识，提升自然语言生成的准确性、可解释性和逻辑推理能力，相对于传统检索增强方法，在医学专业领域可以更好地解决多跳推理不足、语义模糊等问题。

3.2.1 问题定义与核心挑战

将疾病症状与科室精准映射问题形式化定义如下：

输入：

1. 通用疾病知识图谱 $G=(E,R)$ ，其中实体集 E 包含疾病、症状、通用科室等类型，关系集 R 包含“就诊科室”、“典型症状”等。

2. 目标医院 h 的具体挂号科室目录 $L_h=\{d_{h1},d_{h2},\dots,d_{hn}\}$ ，其名称和粒度可能与 G 中的通用科室不一致。

3. 患者的自然语言症状描述 S 。

输出：从列表 L_h 中推荐的匹配度最高的科室 $d_{rec}\in L_h$ ，并可提供简要的推理依据。

核心挑战：直接使用 G 进行匹配会因 G 中的通用科室（如“心血管内科”）与 L_h 中的具体名称（如“心内科门诊”）不匹配而失败。同时，症状描述 S 可能存在模糊性和口语化表达，需要深度的语义理解与临床推理。

3.2.2 方法框架概述

如图 1 所示，本框架依次执行三个模块：① 知识检索与增强：利用症状描述从通用知识图谱中动态抽取相关子图，形成推理的“知识锚点”；② 上下文提示构建：将目标医院科室目录、检索到的相关知识以及任务指令，整合成一个结构化的增强提示；③ 大模型推理与解析：大语言模型基于提示进行推理，输出与目标医院语境对齐的科室推荐。各模块的具体规则见 3.2.3。

整个流程实现了从通用知识到具体医院服务的“语义桥接”。框架整合知识图谱与原始文本，基于逻辑形式直接查询知识图谱，支持结构化推理，减少领域专业不足、事实性错误等问题的出现，能够更好地应对预检分诊中所需的病症到科室的专业领域复杂推理，提高关联及推理的准确率。

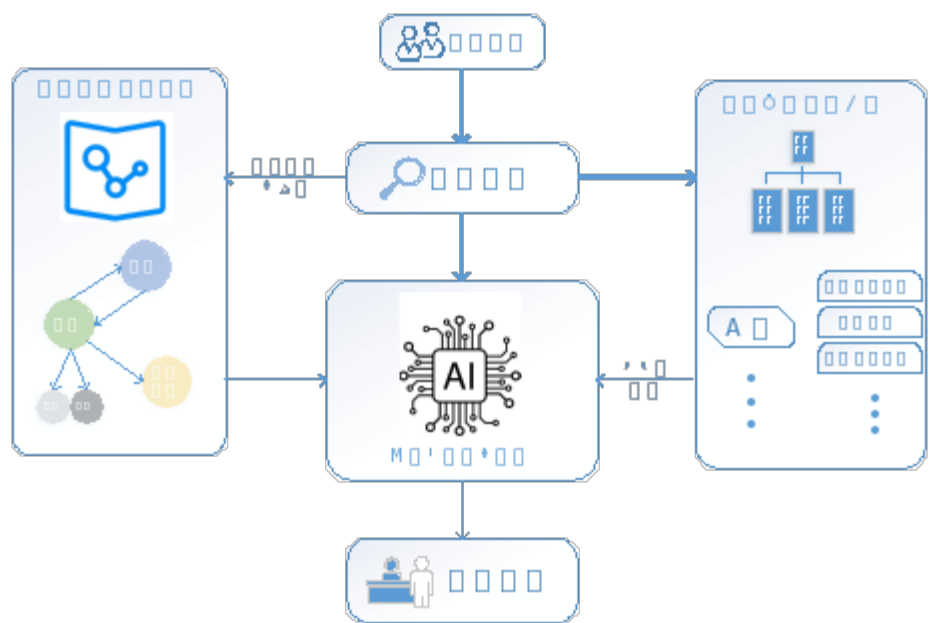


图 1 基于知识增强的大模型预检分诊框架图

3.2.3 知识增强提示构建的具体实现

提示词的质量直接决定了大语言模型的推理性能，同时在本框架中需要大语言模型使各模块串联协作以完成任务。本研究设计了一个多层次、知识增强的提示模板，**主要由四个部分构成**：

系统角色定义部分：固定指令，用于设定大模型的行为模式。
相关知识与上下文注入部分：实现精准映射的关键动态部分，**包含目标医院科室目录及知识图谱子图检索与自然语言化**。目标医院科室目录用于限定模型的输出空间；知识图谱子图检索首先对症状描述 S 进行实体链接和归一化，识别出关键症状实体。随后从通用知识图谱 G 中，以这些症状实体为起点，遍历一跳或两跳关系，抽取包含相关疾病、症状及这些疾病所属的通用科室的子图 G_i 。最后，将 G_i 中的三元组转化为流畅的自然语言陈述，作为背景知识。

任务指令部分：明确任务指令：“请基于上述医学知识，针对以下患者的症状描述，从上述科室目录中选出最应该挂号的科室。请逐步思考，并最终**以JSON格式输出**，包含‘recommended_department’（推荐科室）和‘reasoning’（简要推理步骤）两个字段。”；患者症状描述：“患者症状：<此处插入症状描述 S >。”

输出格式规范部分：要求模型以指定JSON格式输出，便于后续程序化解析与集成，确保结果的稳定性。

其中，跳数策略为：以症状实体为起点，优先抽取一跳关系（症状→疾病、症状→科室）。若一跳返回的疾病实体数 <3 个，或症状描述中含有程度副词（“剧烈”、“持续”、“反复”、“阵发性”），则扩展至两跳（症状→疾病→科室/检查）。禁止两跳以上，以避免噪音引入。子图筛选时，仅保留关系类型属于集合{“典型症状”，“就诊科室”，“鉴别疾病”}的三元组。若存在冲突信息（如同一症状对应多个不同通用科室），本框架将保留所有候选，并将冲突信息作为“鉴别诊断”内容一并输入大语言模型，由模型结合患者具体症状进行综合判断。

通过此提示模板，将目标医院的具体挂号科室系统、通用的医学知识和具体的用户需求整合进一个统一的推理上下文中，极大地提升了大模型在该复杂、细粒度推理映射任务中的准确性与可靠性。

3.2.4 面向异构科室名称的动态语义对齐机制

上述提示构建过程中的“**科室目录作为输出空间约束**”以及“**知识图谱疾病-通用科室关联作为桥接**”共同构成了本研究的动态语义对齐机制。

与传统静态相似度方法（如字符串编辑距离、BERT语义相似度）不同，本机制不预先计算任何相似度矩阵，而是每次根据患者症状和**科室目录**动态生成推理路径。具体地，以通用科室为中间锚点：知识图谱提供的“疾病→通用科室”关系被转换为“疾病→具体医院科室”的候选映射（例如：图谱中“冠心病→心血管内科”，在目标医院 L_h 中可

对应“心内科门诊”或“心脏中心”）；大语言模型综合患者主诉、所有候选映射以及科室目录的整体语义，进行最终选择。

该机制有效解决了名称及结构异构的问题，同时能够应对跨科室、症状模糊等复杂场景（见 4.2.4 节结果分析）。

4 实验与验证

4.1 知识图谱构建

本文基于 MedKaaS 系统，通过筛选网络医学开放数据中的非结构化文本，构建了一个以“症状-疾病-科室”为核心的预检分诊知识图谱，包含 9 类共 9,298 个实体和 16 种共 47,251 条关系（三元组）。实体、关系的详细类型分布如表 1、表 2 所示。

表 1 知识图谱实体统计

类型	子类	数量	占比
实体	疾病	5,727	61.59%
	症状	2,006	21.57%
	其他（科室、检查、药物等）	1,565	16.84%

表 2 知识图谱关系统计

类型	子类	数量	占比
关系	症状→疾病映射	26,939	57.01%
	疾病→科室判别	6,219	13.16%
	其他关系	14,093	29.83%

从实体分布看，疾病与症状实体合计占实体总数的 80% 以上，体现了图谱以临床核心概念为主体。从关系分布看，服务于“症状→疾病”推理的关系占关系总数的 57.01%，直接用于“疾病→科室”判别的关系占 13.16%。该结构充分体现了图谱以分诊任务为中心的设计导向，保证了图谱能够高效支持从患者主诉到科室推荐的核心推理路径。

该图谱为大语言模型提供可信、结构化的医学先验知识，作为大模型的知识基座，辅助和增强大模型对于预检分诊任务中各类医学专业词语的理解，保证框架在智能分诊应用中的可靠性和可用性。

4.2 预检分诊效果

4.2.1 数据集构建

由于缺乏公开的、已标注特定医院科室的症状-挂号对数据集，本研究采用“大模型模拟生成+专家人工校验”的策略构建评估数据集。

模拟数据生成：以自建的预检分诊知识图谱为基础，选取 30 余种常见疾病，覆盖内科、外科、眼科等 18 个科室类别。针对每种疾病，使用大语言模型（DeepSeek-R1-Lite-Preview）生成 3-5 条不同表述方式和症状组合方式的患者主诉，共计 135 条原始数据。

数据清洗与审核：首先检测数据是否存在错漏或科室不匹配等情况，确认子集标签与科室标签内容的正确性（以北京协和医院东单院区科室目录为准）。不一致意见由审核者确认。最终剔除 6 条主诉简单或与其他主诉重复的数据，修正 11 条子集与科室标签，形成 129 条有效数据。

剔除低样本科室：为确保训练/测试集划分的有效性，剔除仅有 1 个样本的 3 个科室对应的主诉数据（共 3 条）。剔除后剩余 126 条有效数据。

“复杂案例”子集构建：为确保评估全面性，从 126 条数据中筛选出 31 条“复杂案例”。筛选标准为满足以下任一条件：(a) 症状描述涉及两个及以上可能科室（如“胸痛伴发热”可对应心内科或呼吸科）；(b) 症状描述模糊或非典型（如“肚子不舒服”），容易与其他疾病混淆；(c) 对应疾病为罕见病或跨科室疾病（如食管胃底静脉曲张破裂出血涉及消化内科及血管外科）。该子集用于评估模型在困难场景下的鲁棒性。

数据集划分与样本分布：在基线模型训练与测试中，将 126 条数据按照 8: 2 的比例随机划分，使用分层抽样保证各科室类别在训练/测试集中的比例与总体分布一致。最终分为训练集（100 条）和测试集（26 条）。其中样本分布占比前三的科室为：心内科门诊（19.38%），基本外科门诊（15.5%）和消化内科门诊（13.18%）。

10.12201/bmr.202608.00014V1

4.2.2 评估指标与实验设置

评估指标：采用 Top-1 准确率（ACC@1）和 Top-3 准确率（ACC@3）作为主要评估指标。同时，为评估模型在复杂场景下的性能，在“复杂案例”子集上额外计算宏平均 F1 分数。

实验设置：本方法及所有涉及大模型的基线均采用相同的大模型服务与接口（DeepSeek-R1-Lite-Preview）。

4.2.3 主实验与消融实验

主实验：对比所提 KGE-LLM TF 方法、基线模型（BERT、Sentence-BERT、TransE）、与无知识增强 LLM 方法的性能。其中 BERT、Sentence-BERT 使用数据集划分的训练集和测试集进行测试；TransE 方法与文本相同的知识图谱进行训练；无知识增强 LLM 方法为不引入知识图谱，仅将症状描述和目标医院科室目录输入大模型，直接推荐科室。

消融实验：为剖析各组件贡献，设置以下消融变体：w/o KG 为去除知识图谱子图检索与增强模块；w/o CoT 为移除“逐步思考”的思维链指令。

4.2.4 结果分析与讨论

主实验结果如表 3 所示。所提出的 KGE-LLM TF 方法在三个评价指标（ACC@1、ACC@3、复杂案例 F1）上均取得最优结果，尤其在 ACC@1（84.91%）与复杂案例 F1（75.71%）上显著优于其他基线模型，说明融合知识图谱增强与思维链推理能有效提升模型在复杂场景下的判别能力。

传统表示学习模型（BERT、Sentence-BERT、TransE）在 ACC@1 与复杂案例 F1 上表现相对有限。虽然 Sentence-BERT 在 ACC@3 上与 KGE-LLE TF 相差不大，但在复杂案例 F1 上最高仅为 50.00%，说明其对深层语义与复杂推理的建模能力不足。ACC@1 的表现不佳也说明模型算法不能稳定推荐最有可能的结果，增加了使用者的判断成本。

相较于无知识增强 LLM，本研究的方法在 ACC@1 上提升了 5.67%，这有力证明了外部知识图谱对于约束大模型“幻觉”、提供可靠推理依据的关键作用。同时，本方法在“复杂案例”子集上的 F1 分数高出 30% 左右，说明在症状描述模糊、涉及多科室的复杂推理任务中，“知识锚点”的存在能为大模型的推理能力提供基点和帮助。

消融实验结果表明，知识图谱增强和思维链提示均是有效组件。移除知识图谱（w/o KG）导致性能下降最为明显，特别是在 ACC@1 指标上下降 18.87%，这表明动态检索的相关知识是精准映射的基础；复杂案例 F1 下降至 38.89%，说明知识增强对复杂推理具有关键作用。移除思维链（w/o CoT）在 ACC@1 和复杂案例 F1 上分别下降约 9.44% 和 16.19%，表明分步推理机制能有效提升模型决策的可解释性与准确性，尤其在处理复杂案例时效果突出。

表 3：主实验与消融实验结果（主要指标，%）

方法	ACC@1	ACC@3	复杂案例 F1
BERT	19.23	50.00	16.67
Sentence-BERT	50.00	92.31	50.00
TransE	27.91	51.16	14.88
无知识增强 LLM	79.24	88.68	45.83
KGE-LLM TF (Ours)	84.91	92.45	75.71
消融 w/o KG	66.04	84.91	38.89
消融 w/o CoT	75.47	86.79	59.52

整体实验证明，将大模型的开放域泛化能力与知识图谱的结构化、可信知识相结合，是实现可靠医疗决策支持的有效范式。本框架的核心优势在于其“即插即用”的适应性，适配一家新医院，仅需输入其科室目录，无需重新训练模型，极大降低了部署成本，为知识图谱的快速落地提供了可行路径。

在推理依据一致性方面，使用本研究框架模型对选取的 32 条测试样本独立的生成 3 次包含推理依据的结果，并使用 Sentence-BERT 算法（模型为 sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2）计算语义相似度，平均语义相似度为 0.8554，标准差为 0.0554。由于未对推理依据做出严格约束，模型生成描述相差较大，但整体都呈现出主诉分析、外部知识库推理、科室对应三个推理阶段如在某主诉的分析结果中，有的推理以患者病情记录（“青年男性，急性起病，先有发热肌痛、关节痛等症状，后出现胸闷、呼吸困难、心悸、心律不齐…”）的方式，有的推理

10.12201/bmr.202608.00014V1

以症状罗列（“关键症状‘发热、关节酸痛后出现胸闷、气短、心悸…’”）的方式呈现但其主要内容仍具有一致性。综上，本文提出的 KGE-LLM TF 方法通过知识图谱增强与思维链推理的协同作用，显著提升了模型在复杂场景下的性能，消融实验进一步验证了两个模块的有效性与必要性。但同时模型框架受限于知识图谱的覆盖范围与时效性，对于新出现的疾病或亚专科，分诊推理可能失效，尤其当知识图谱中相关疾病与病症关系不够明确时，准确率可能降低。另一方面，基于大模型的不确定性与上下文长度限制，模型的输出可能出现随机性波动，当目标医院科室目录 L_h 极长时，完整的列表可能与复杂的知识子图一同耗尽模型的上下文窗口。

5 结语

在医疗数字化转型与人工智能技术迭代的双重背景下，构建可信赖的临床决策支持系统（CDSS）是智慧医院建设的核心诉求。本研究提出的 KGE-LLM TF 框架，通过引入外部知识增强与动态推理机制，有效解决了异构科室环境下的语义对齐难题。实验结果表明该框架在处理复杂、非典型症状的映射任务时，性能显著优于基线模型及无增强模型（提升幅度>25%），ACC@1 与 ACC@3 分别达到 84.91%与 92.45%。这一结果不仅验证了框架的有效性，更为解决大模型在垂直领域落地的“幻觉”与“适配”两大痛点提供了实证参考。此外，本研究不仅关注算法精度，更强调工程可行性与部署效率。所提方案具备良好的跨机构迁移能力，无需对医院现有数据结构进行重构，具有极高的推广价值。当然，本研究在在多学科交叉验证的深度上仍有拓展空间，未来可探索该框架对多模态数据（如体征监测数据）的融合、危急重症的分级预警机制，以及更高效的知识检索与提示压缩技术，以进一步释放 AI 技术在实时诊疗、基层医疗、远程诊疗等临床场景中的普惠价值。

作者贡献：刘会负责研究设计、文献调研、论文撰写；姚宽达、刘湘闽负责方法设计、实验验证、图表绘制；胡佳慧、王茜负责研究选题、论文修订；方安负责提供指导、论文审阅与修订。
利益声明：所有作者均声明不存在利益冲突。

参考文献

[1] 刘瑶. 优质护理服务用于门诊预检分诊中的效果[J]. 生命科学仪器,2023,21(z1):209.DOI:10.11967/2023006195.
[2] 牟丽虹.浅谈预检分诊的体会[J].健康必读,2020,(22):262.
[3] 刘徐方,刘瑞雪.智能化预检分诊系统在急诊患者中的应用效果[J].中国民康医学,2024,36(02):190-192.
[4] 苗开贵,沈亚奇,王沛沛.智能化急诊预检分诊系统对急诊预检分诊质量的影响[J].齐鲁护理杂志,2022,28(20):160-162.
[5] 朱超宇,刘雷.基于知识图谱的医学决策支持应用综述[J].数据分析与知识发现,2020,4(12):26-32.
[6] 胡红娟,周阳,匡泽民,等.医学知识图谱应用研究进展[J].医学信息学杂志,2022,43(5):30-33, 39
[7] 刘悦悦,李燕.医学知识图谱研究综述[J].软件导刊, 2023, 22 (5) : 241-247.
[8] Gene Ontology Consortium. Gene Ontology Consortium: going forward. Nucleic Acids Res. 2015 Jan;43(Database issue):D1049-56. doi: 10.1093/nar/gku1179. Epub 2014 Nov 26. PMID: 25428369; PMCID: PMC4383973.
[9] 郭琳,陈晓慧,肖梅.知识图谱研究综述[J].信息记录材料, 2023, 24(6):17-19.
[10] Wishart DS, Feunang YD, Guo AC, Lo EJ, Marcu A, Grant JR, Sajed T, Johnson D, Li C, Sayeeda Z, Assempour N, Iynkkaran I, Liu Y, Maciejewski A, Gale N, Wilson A, Chin L, Cummings R, Le D, Pon A, Knox C, Wilson M. DrugBank 5.0: a major update to the DrugBank database for 2018. Nucleic Acids Res. 2018 Jan 4;46(D1):D1074-D1082. doi: 10.1093/nar/gkx1037. PMID: 29126136; PMCID: PMC5753335.
[11] 奥德玛,杨云飞,穗志方,等.中文医学知识图谱 CMeKG 构建初探[J].中文信息学报,2019,33(10):1-9.
[12] 闻龙,卢若谷,种璟,罗之沁,吴娜.医学知识图谱构建与应用的研究[J].长江信息通信,2023,36(10):1-8.DOI:10.3969/j.issn.1673-1131.2023.10.001.
[13] 蒋君,肖宇锋,门佩璇,等.医疗健康知识图谱需求与技术构建研究[J].医学信息学杂志,2024,45(3):40-45

- [14] 刘燕,张潇潇,侯丽.面向医药卫生知识服务系统的学术知识图谱构建与应用研究[J].医学信息,2024,45(4):1-7,30.DOI:10.3969/j.issn.1673-6036.2024.04.001.
- [15] 赵继贵,钱育蓉,王魁,等.中文命名实体识别研究综述[J].计算机工程与应用,2024,60(01):15-27.
- [16] Ge J, Sun S, Owens J, Galvez V, Gologorskaya O, Lai JC, Pletcher MJ, Lai K. Development of a liver disease-specific large language model chat interface using retrieval-augmented generation. *Hepatology*. 2024 Nov 1;80(5):1158-1168. doi: 10.1097/HEP.0000000000000834. Epub 2024 Mar 7. PMID: 38451962; PMCID: PMC11706764.
- [17] Fink A, Rau A, Reisert M, Bamberg F, Russe MF. Retrieval-Augmented Generation with Large Language Models in Radiology: From Theory to Practice. *Radiol Artif Intell*. 2025 Jul;7(4):e240790. doi: 10.1148/ryai.240790. PMID: 40464682.
- [18] Miao J, Thongprayoon C, Suppadungsuk S, Garcia Valencia OA, Cheungpasitporn W. Integrating Retrieval-Augmented Generation with Large Language Models in Nephrology: Advancing Practical Applications. *Medicina (Kaunas)*. 2024 Mar 8;60(3):445. doi: 10.3390/medicina60030445. PMID: 38541171; PMCID: PMC10972059.
- [19] Wang S, Yang H, Liu W. Research on the construction and application of retrieval enhanced generation (RAG) model based on knowledge graph. *Sci Rep*. 2025 Nov 18;15(1):40425. doi: 10.1038/s41598-025-21222-z. PMID: 41253865; PMCID: PMC12627698.
- [20] Ye C, Swiers R, Bonner S, Barrett I. A Knowledge Graph-Enhanced Tensor Factorisation Model for Discovering Drug Targets. *IEEE/ACM Trans Comput Biol Bioinform*. 2022 Nov-Dec;19(6):3070-3080. doi: 10.1109/TCBB.2022.3197320. Epub 2022 Dec 8. PMID: 35939454.
- [21] Feng Y, Zhou L, Ma C, Zheng Y, He R, Li Y. Knowledge graph-based thought: a knowledge graph-enhanced LLM framework for pan-cancer question answering. *Gigascience*. 2025 Jan 6;14:giae082. doi: 10.1093/gigascience/giae082. PMID: 39775838; PMCID: PMC11702363.
- [22] Xia Y, Li J, Deng B, Huang Q, Cai F, Xie Y, Sun X, Shi Q, Dan W, Zhan Y, Jiang L. Knowledge Graph-Enhanced Deep Learning Model (H-SYSTEM) for Hypertensive Intracerebral Hemorrhage: Model Development and Validation. *J Med Internet Res*. 2025 Jun 12;27:e66055. doi: 10.2196/66055. PMID: 40505141; PMCID: PMC12203281.
- [23] 王博,于志昊,张军雁,等.基于电子病历数据和知识增强的医疗大语言模型构建方法研究[J].解放军医学院学报,2025,46(01):97-103+119.
- [24] Bordes A, Usunier N, Garcia-Duran A, et al. Translating Embeddings for Modeling Multi-relational Data[J]. Curran Associates Inc. 2013.
- [25] Sun Z, Deng Z H, Nie J Y, et al. RotatE: Knowledge Graph Embedding by Relational Rotation in Complex Space[J]. 2019.DOI:10.48550/arXiv.1902.10197.
- [26] Devlin J, Chang M W, Lee K, et al. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding[J]. ArXiv, 2019, abs/1810.04805.DOI:10.18653/v1/N19-1423.
- [27] J. Li, A. Sun, J. Han and C. Li, "A Survey on Deep Learning for Named Entity Recognition," in *IEEE Transactions on Knowledge and Data Engineering*, vol. 34, no. 1, pp. 50-70, 1 Jan. 2022, doi: 10.1109/TKDE.2020.2981314.
- [28] Reimers N, Gurevych I. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks[J]. 2019.DOI:10.18653/v1/D19-1410.