

融合外部知识和提示学习的电子病历命名 实体识别方法

李阳¹，曹子佳¹，肖万里¹，赵文涵¹

(¹扬州大学附属医院 扬州 225000)

【摘要】 **目的/意义** 提升对电子病历中医学实体的识别能力，为临床决策支持与医学知识库构建提供基础支撑，助力医院信息化建设。**方法/过程** 提出融合外部知识和提示学习的电子病历命名实体识别模型，首先，利用类别描述文本初始化可学习领域提示向量，与原文本拼接后通过 ERNIE 模型生成聚焦领域特征的词向量；其次，使用 softLexicon 方法和注意力机制融合外部词典，增强对模糊实体与边界的表征能力；最后，使用两阶段识别机制，通过候选跨度剪枝与困难负样本生成提升边界检测精度，并利用注意力机制将候选跨度与提示向量交互以实现分类。**结果/结论** 在四个公开电子病历数据集上进行对比实验和消融实验，实验结果证明了模型的有效性。

【关键词】 命名实体识别；提示学习；融合外部知识；两阶段识别机制

【文献标志码】 A

Named Entity Recognition model for Electronic Medical Records Based on External Knowledge Integration and Prompt Learning

LI Yang¹, CAO Zijia¹, XIAO Wangli¹, ZHAO Wenhan¹

¹Affiliated Hospital of Yangzhou University, Yangzhou 225000, China

【Abstract】 **Purpose/Significance** To enhance the recognition capability of medical entities in electronic medical records, provide foundational support for clinical decision support and medical knowledge base construction, and facilitate hospital informatization. **Method/Process** named entity recognition model for electronic medical records based on external knowledge integration and prompt learning is proposed. First, learnable domain-specific prompt vectors initialized from category description texts are constructed and concatenated with the original text to generate domain-focused word embeddings via the ERNIE model. Second, the softLexicon method and an attention mechanism are employed to integrate external dictionary, enhancing the model's ability to represent ambiguous entities and boundaries. Finally, a two-stage recognition mechanism is adopted, where candidate span pruning and hard negative sample generation improve boundary detection accuracy, and an attention mechanism is used to enable interaction between candidate spans and prompt vectors for classification. **Results/Conclusion** Comparative experiments and ablation studies were conducted on four publicly available electronic medical record datasets, and the experimental results demonstrate the effectiveness of the proposed model.

【Key words】 named entity recognition; prompt learning; external knowledge integration; two-stage recognition mechanism;

【作者简介】李阳，助理工程师，发表论文 3 篇，主要研究方向为医学自然语言处理，医疗数据挖掘；

通讯作者：曹子佳，高级工程师。

【基金项目】江苏省医院协会医院管理创新研究课题 JSYGY-3-2024-435。

10.12201/bmr.202608.00018V1

1 引言

命名实体识别旨在从非结构化文本中识别并分类实体片段，是自然语言处理的基础任务^[1-2]。在医学领域，随着《“健康中国2030”规划纲要》和《“十四五”全民健康信息化规划》的深入推进，电子病历信息化建设已成为医院高质量发展的重要任务。从电子病历中自动提取疾病、症状、检查、药物等医学实体^[3-4]，对于构建医学知识图谱^[5-6]、提升临床决策支持系统的智能化水平^[7]及研发智能问诊系统^[8-9]具有重要的应用价值。

然而，电子病历实体识别面临三大挑战：第一，医学术语语义高度依赖上下文，导致实体识别困难；第二，高质量标注语料构建成本高昂；第三，句法复杂且实体嵌套现象普遍，边界检测与类别判定难度大。

早期医学领域的命名实体识别方法主要分为基于规则和基于统计机器学习两类，前者解释性强但泛化能力弱^[10]，后者虽性能较好但依赖人工特征^[11]。随着深度学习兴起，Liu等^[12]提出BiLSTM-CRF模型，利用双向长短期记忆网络捕获上下文信息并用条件随机场确保标签序列全局最优，Devlin等^[13]提出预训练语言模型BERT，极大提升了实体识别性能，Liu等^[14]进一步利用预训练模型提取全局与多层次语义特征。近年来，外部知识融合与提示学习成为研究热点，Ma等^[15]提出softLexicon方法将外部字典词汇信息与字符信息融合，He等^[16]通过提示将词汇和词典信息注入BERT增强细粒度语义关联，Mai等^[17]针对少样本中文命名实体识别提出了基于分段的提示调优方法。

针对上述挑战，本文提出融合外部知

识与提示学习的电子病历命名实体识别模型。核心思想是：首先，为每个实体类别构建可学习领域提示向量，引导模型聚焦医学语义；其次，融合外部医学词典知识，结合BiLSTM增强上下文依赖，提升对模糊实体及边界的识别能力；最后，使用“先边界、后类别”的两阶段机制处理嵌套实体，通过候选跨度剪枝与困难负样本生成提升边界检测精度，并利用注意力机制实现精确分类。

2 模型结构

本节阐述模型的整体架构。如图1所示，包含四个核心模块：提示学习模块使用可学习领域提示向量引导模型聚焦医学领域特征；融合外部知识模块融合ERNIE、外部词典和上下文依赖，生成语义丰富的词向量；实体边界检测模块通过枚举、筛选判定候选文本跨度是否为实体；实体类别识别模块对检测出的实体跨度进行精确分类。

2.1 提示学习模块

本模块依据数据集使用手册的类别描述文本为每个实体类别（含“非实体”类别）构建可学习领域提示向量，如症状的描述文本为“疾病所引起患者的主观感受或客观征象”。将描述文本初始化后的提示向量与原文本拼接，共同输入ERNIE模型，以引导模型更加关注医学领域特征，从而缓解领域差异与数据稀缺问题。

2.2 融合外部知识模块

为充分利用医学领域的词典知识，本模块在ERNIE输出层后引入softLexicon方法^[15]，并通过注意力机制融合外部词典知识与上下文表征，再经BiLSTM增强序列建模能力，流程如图2所示。softLexicon方法为输入文本每个字符匹配外部词典词汇，若匹配词汇在输入文本中存在则为有效词汇，最终将检索结果分为四个集合，分别是以该字符开头的集合 B ，包含该字符但非首尾的集合 M ，以该字符结尾的集合 E 和单独成词的集合 S 。如图2中字符“肠”，四个集合分别是 $B=\{\text{肠癌}\}$ 、 $M=\{\text{直肠癌}\}$ 、 $E=\{\text{低位直肠, 直肠}\}$ 和 $S=\{\text{肠}\}$ 。根据外部词向量集将词汇映射为词向量，并取各集合中词向量均值作为集合表示 V_i 。

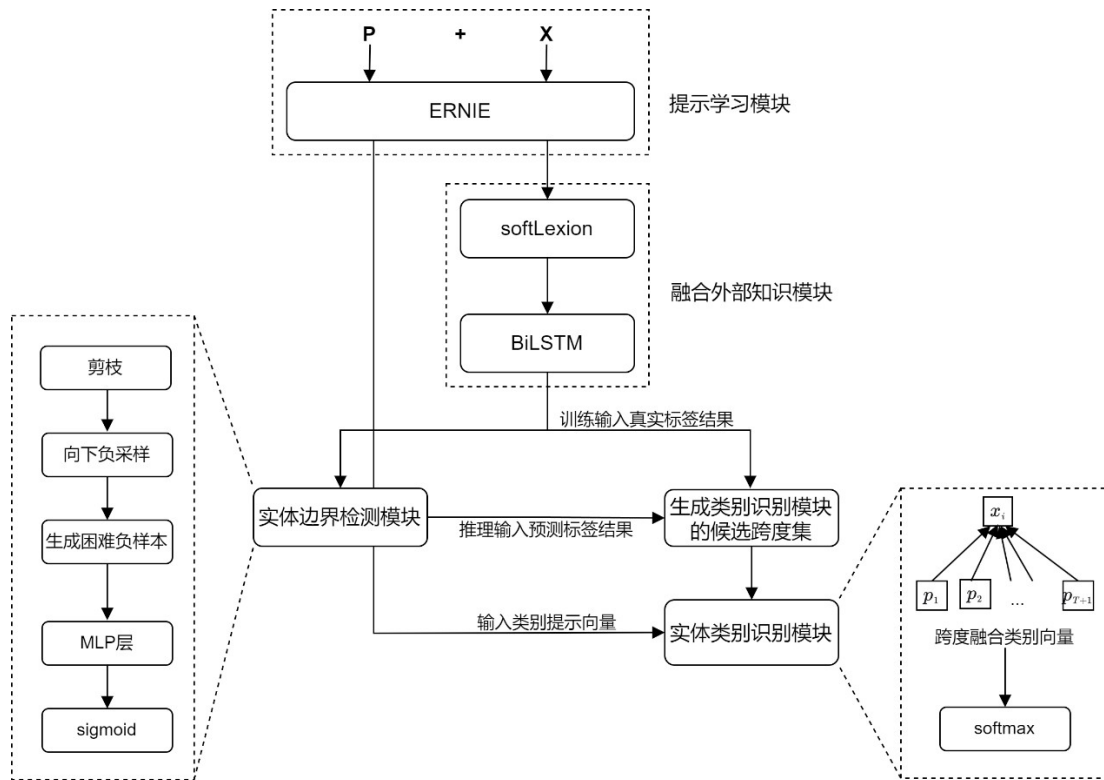


图 1 模型整体框架图

需要指出的是，softLexicon 的效果受外部词典覆盖范围与质量制约，作为静态资源难以完全覆盖电子病历中动态多样的临床表述，目标语料中的实体表述与词典词汇有偏差时，可能匹配到与上下文无关的词汇而引入噪声。为此，本文使用注意力进行动态加权融合：不将四个集合等权平均，而根据当前字符的上下文语义自适应计算各集合贡献权重，使噪声集合权重趋近零，与当前语境高度相关的词汇集合获得较大权重。这种机制能在保留有效词典知识的同时自适应缓解噪声，增强对实体边界与类别语义的鲁棒性。具体来说，使用注意力计算每个字符与四个词汇集合向量的交互权重，生成融合后的词汇增强向量 u_i 。因为字符表示 h_i^{ERNIE} 和 u_i 共享同一语义空间，故将两者直接相加，得到融合外部知识的字符表示 h_i^{fuse} 。

计算公式如下：

$$V = (h_i^{ERNIE} W_{att} \tilde{v})^T \quad (1)$$

$$A_i = \text{softmax}(V) \quad (2)$$

$$u_i = A_i V_i \quad (3)$$

$$h_i^{fuse} = h_i^{ERNIE} + u_i \quad (4)$$

为增强对序列方向性及长距离语义依赖的建模能力，本文将融合外部词典的向量 H^{fuse} 送入 BiLSTM：

$$H^{LSTM} = \text{BiLSTM}(H^{fuse}) \quad (4)$$

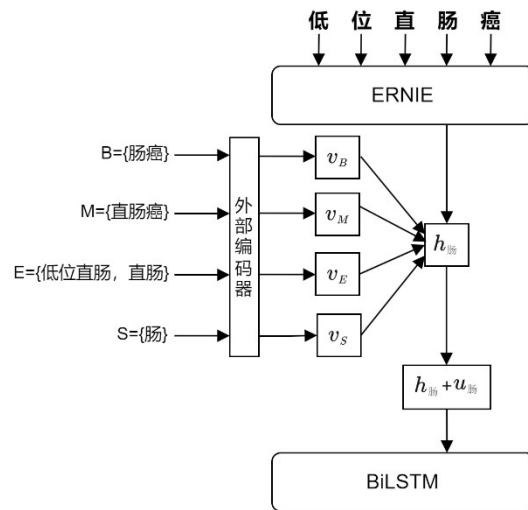


图 2 外部知识增强的编码器模块

2.3 实体边界检测模块

为处理嵌套实体问题，本文使用两阶段识别机制，该模块旨在从电子病历文本中检测实体的起止边界。首先，在预定义的最大跨度长度 L 内枚举所有候选跨度。鉴于许多跨度明显不符合实体构成规则，本文提出基于先验知识的剪枝策略过滤低置信度候选跨度：剔除以停用词（如“的”、“了”）或标点为头尾的跨度；剔除包含不匹配标点对（如只有左括号无右括号）的跨度；剔除仅由数字或单个非中文字符构成的跨度。该策略能减少候选集规模，并使模型关注真实实体。

剪枝后，负样本仍远多于正样本。为此，本文在每轮训练中进行向下负采样，将正负样本比例固定为 1:5。为进一步提升

10.12201/bmr.202608.00018V1

模型对模糊实体边界的判别力，本文使用困难负样本生成机制。困难负样本定义为：与真实实体在文本位置上重叠，但起止边界不完全一致的非实体跨度，这类样本与正例语义相似度高，仅边界存在细微偏差。生成方式为：对每个真实实体随机偏移其起始或结束位置（左移或右移 1-3 个字符），同时确保新跨度与原实体有重叠但不重合。如真实实体“颅骨缺损”，可生成“骨缺损”或“颅骨缺”作为困难负样本。通过引入这类边界高度相似的样本，迫使模型学习更精细的边界判别特征，从而提升对复杂文本边界的鲁棒性。为处理样本不平衡问题，边界检测模块使用 Focal Loss 作为损失函数。

2.4 实体类别识别模块

该模块对边界检测模块输出的候选实体跨度进行细粒度分类，为充分利用类别语义信息，本模块引入注意力机制，使候选跨度向量 h_i^{span} 与所有提示向量 H^{prompt} 交互，得到类别感知的增强表示 z_i ，最后将 z_i 和跨度向量 h_i^{span} 拼接，计算公式如下：

$$h_i^{span} W_{attn2}(H^{prompt}; T) \\ A_i^{diff} = \text{softmax}(\dot{A}) \quad (5)$$

$$z_i = A_i^{diff} H^{prompt} \quad (6)$$

$$z'_i = w[h_i^{span}; z_i] + b \quad (7)$$

训练时使用真实标签输入类别识别模块，避免误差传播；推理时使用预测的候选跨度。模型通过加权求和两个模块的损失进行端到端联合训练。

3 实验设计与结果分析

3.1 实验数据集及实验设置

本文使用四个公开数据集：CCKS2017、CCKS2019、CMeEE 和 DiaKG，均基于临床电子病历标注，其中 CCKS 和 CMeEE 包含多种医学实体，DiaKG 专注于糖尿病领域。使用过程中严格遵循医疗数据隐私保护规范。统计信息如表 1 所示。

表 1 实验数据集

数据	实体类型	标注实体数	嵌套实体数	任务类型
CCKS2017	5	31865	-	扁平实体

CCKS2019	6	23020	-	扁平实体
CMeEE	9	25098	990	嵌套实体
DiaKG	18	18094	2474	嵌套实体

本文实验环境配置如下：深度学习框架为 PyTorch。ERNIE 模型为 ERNIE-Health-zh，是百度基于 ERNIE 打造的中文医疗领域预训练语言模型，通过医疗实体掩码和医疗问答匹配等任务在中文医学语料上预训练，能更精准地捕捉中文医学语义，包含 12 层 Transformer 编码器，隐藏层维度为 768。优化器为 AdamW。外部词典为《常用临床医学名词（2023 年版）》，外部词向量集为腾讯 AI Lab 发布的预训练中文词向量集，其余关键超参数如表 2 所示。

表 2 超参数设定

超参数	设定值	超参数	设定值
批次大小	8	Dropout	0.1
学习率	3e-5	权重衰减	0.01
训练轮次	25	最大文本长度	512
最大跨度长度	30	ρ (困难负样本比例)	0.5

3.2 评价指标

本文使用准确率（Precision）、召回率（Recall）与 F1 值（F1-score）进行评估。其中，准确率为模型预测为正例的实体中真正为正例的比例；召回率为所有真实正例中被模型正确预测的比例；F1 值为准确率与召回率的调和平均数，用于综合反映模型在两项指标上的整体性能。计算公式如下：

$$F1 = \frac{2PR}{P+R} \times 100 \quad (8)$$

3.3 对比实验

为验证模型的有效性，本文在四个数据集上开展了对比实验。所选对比模型包括序列标注模型（BiLSTM+CRF、Lattice+LSTM、FLAT^[18]、MUSA^[19]和 M-DSCMG^[20]）、嵌套实体识别模型（Global Pointer^[21]、W2NER^[22]和 RoBERTa-wwm-ext-GCA-crf^[23]）以及通义千问团队发布的轻量级大语言模型 Qwen3-0.6B。同时，为验证医学预训练语言模型的有效性，额外增设了 BERT 作为对照组。

表 3 CCKS2017 和 CCKS2019 的实验结果（%）

模型	CCKS2017			CCKS2019		
	P	R	F1	P	R	F1
BiLSTM+CRF	91.17	86.00	88.12	81.49	80.52	81.00
Lattice+LSTM	90.04	89.01	89.52	80.56	81.47	81.08
FLAT	90.31	92.72	91.51	82.16	82.57	82.36
MUSA	92.67	90.97	91.81	83.31	82.44	82.87
M-DSCMG	92.36	92.24	92.30	83.82	87.27	85.51
Qwen3-0.6B	92.57	91.54	92.05	84.39	87.45	85.89
Ours(BERT)	93.42	92.94	93.18	86.51	88.37	87.43
Ours(ERNIE)	94.24	93.65	93.95	87.29	89.18	88.22

表 4 CMeEE 和 DiaKG 的实验结果 (%)

模型	CMeEE			DiaKG		
	P	R	F1	P	R	F1
Global Pointer	60.61	62.81	61.69	84.27	83.46	83.79
W2NER	65.69	60.09	62.77	83.56	84.20	83.87
RoBERTa-www-ext-GCA-crf	66.43	63.07	64.36	87.24	83.54	85.35
Qwen3-0.6B	63.68	62.48	63.07	85.12	83.77	84.44
Ours(BERT)	64.38	65.88	65.12	87.34	86.02	86.68
Ours(ERNIE)	65.39	66.92	66.14	88.54	86.82	87.67

实验结果如表 3 和表 4 所示，本文模型取得最佳性能，F1 值分别为 93.95%、88.22%、66.14% 和 87.67%。这是由于本文基于中文医学领域预训练语言模型，使用类别描述文本初始化可学习领域提示向量，引导模型聚焦于医学语义特征。编码阶段使用 softLexicon 构建外部词汇集合并用注意力融合，帮助识别模糊实体和边界；解码阶段利用注意力使候选跨度与提示向量交互，提高分类精度。ERNIE-Health-zh 凭借其在医学语料上的深度预训练，获得比 BERT 更好的结果。而 Qwen3-0.6B 虽具备高效推理能力，但仍低于最佳性能，因为其缺乏医学领域深度适配，且参数量较小，在细粒度医学实体边界和嵌套结构识别上不如专用架构。

在 CCKS2017 和 CCKS2019 数据集的其他对比模型中，M-DSCMG 使用多种注意力融合字符与词汇特征并增强模型识别医学实体边界的能力。MUSA 使用多头自注意力融合外部知识，FLAT 提出扁平化网络结构融合字符和词汇特征。BiLSTM+CRF 和 Lattice+LSTM 是序列标注基线模型。

在 CMeEE 和 DiaKG 数据集的其他对比模型中，RoBERTa-wwm-ext-GCA-crf 使用更复杂的 RoBERTa 生成词向量，并利用循环神经网络和多头注意力学习深层语义知识。W2NER 将词对的相邻关系和头尾关系转换为二维关系表，并通过二维卷积捕捉语义关联。Global Pointer 使用乘法注

意力融合相对位置信息，并识别实体头部和尾部解决训练和推理不一致的问题。

3.4 消融实验

3.4.1 领域提示向量的影响 为验证领域提示向量的有效性，设置无提示和无跨度类别注意力作为对照组。结果如表 5 所示，无提示时模型在各数据集上的 F1 值均最低，表明缺乏领域引导时模型难以准确捕捉医学实体。使用提示虽有提升，但在分类时跨度向量未能充分融合提示向量语义信息。相比之下，本文模型在四个数据集上均取得最优性能。该结果表明，自适应提示能有效引导预训练语言模型关注领域特定特征，跨度类别注意力能有效融合提示向量语义信息，从而提升对复杂医学术语与嵌套结构的识别能力。

3.4.2 融合外部词典对模型的影响 为验证融合外部词典对模型的影响，设置未融合外部词典和平均融合外部词典作为对照组，其中平均融合即将四个词汇集的注意力权重设为相同值。结果如表 6 所示，未融合外部词典在各数据集上的 F1 值均最低，表明外部词典知识能够为字符表示引入词汇边界信息和医学语义先验，在应对复杂实体和模糊边界时发挥作用。平均融合虽引入了外部知识，但四个集合中的噪声词汇与有效词汇被平等地注入词汇向量，导致错误传播。相比之下，本文模型在四个数据集上均取得最优性能。该结果表明，通过注意力可以使与当前上下文相关性较大的词汇集合获得更高的融合权重，缓解无关或错误匹配词汇的噪声干扰。

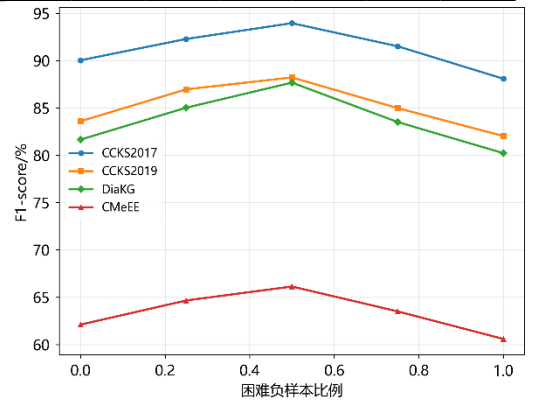
表 5 领域提示向量的影响 (%)

模型	CCKS2017			CCKS2019			CMeEE			DiaKG		
	P	R	F1	P	R	F1	P	R	F1	P	R	F1
无提示	91.7	90.7	91.2	82.8	81.7	82.2	61.9	58.8	60.3	82.7	81.0	81.8
无跨度类别注意力	4	7	5	1	1	6	3	7	6	3	2	6
	92.2	92.4	92.3	86.8	83.9	85.3	64.11	63.4	63.7	86.9	83.1	85.0
Ours	6	0	3	2	1	4	6	8	2	9	1	1
	94.2	93.6	93.9	87.2	89.1	88.2	65.3	66.9	66.1	88.5	86.8	87.6
	4	5	5	9	8	2	9	2	4	4	2	7

表 6 融合外部词典的影响 (%)

模型	CCKS2017			CCKS2019			CMeEE			DiaKG		
	P	R	F1	P	R	F1	P	R	F1	P	R	F1
未融合外部词典	92.8	91.1	91.9	87.5	83.8	85.6	65.4	61.0	63.1	84.6	84.9	84.78
平均融合外部词典	5	0	7	9	8	9	2	7	3	3	3	86.25
	93.4	92.1	92.7	87.4	86.8	87.1	65.4	64.0	64.7	86.6	85.8	86.25
Ours	2	8	9	4	2	3	3	7	4	2	8	87.67
	94.2	93.6	93.9	87.2	89.1	88.2	65.3	66.9	66.1	88.5	86.8	87.67
	4	5	5	9	8	2	9	2	4	4	2	

3.4.3 困难负样本比例的影响 为探究困难负样本比例对模型的影响，本文在 0.0~1.0 范围内设置 5 个比例进行对比实验。结果如图 3 所示，当困难负样本比例为 0.5 时，模型在四个数据集上均取得最优性能。若比例过低，模型对边界模糊样本的识别能力不足；比例过高则易使模型过度拟合困难负样本，导致对真实实体的识别精度下降。0.5 在增强边界判别能力与维持模型泛化性能之间达到最佳平衡，说明适度的困难样本能有效提升模型对复杂文本结构的区分度。



10.12201/bmr.202608.00018V1

图3 四个数据集不同困难负样本比例的F1值

3.5 模型性能及错误分析

为展示模型对数据集中各类别的识别性能，本文选取 CCKS2019 和 CMeEE 展示实验结果，结果如表 7 和表 8 所示。模型在两个数据集上 F1 值分别达到 88.22% 和 66.14%，在 CCKS2019 数据集中，药物的 F1 值为 93.94%，但手术和检验的 F1 值分别为 60%和 80.75%。在 CMeEE 数据集中，疾病、药物和微生物类的 F1 值分别为 77.64%、75%和 73.68%，但检查和检验、科室的 F1 值分别为 38.2%和 57.14%。

实体识别错误的原因包括：数据集中存在不规则长实体，如手术实体“剖腹探查+膀胱旁肿物切除+肠表面肿物切除术”使用“+”作为连接符且长度较长，模型难以准确检测其边界。数据集中的样本分布不均衡，如 CCKS2019 中的检验实体和 CMeEE 中的科室实体在训练集样本中的占比分别仅为 4.46%和 1.44%；数据量过少导致模型对这类实体泛化能力不足。数据集中包含英文简称实体，如检查和检验实体“肌酸激酶同工酶（CK-MB）”，模型难以通过语义特征判别实体类别，因此未能识别出“CK-MB”。

表7 CCKS2019 各个实体类别的实验结果（%）

类型	CCKS2019		
	P	R	F1
解剖部位	89.20	88.90	89.05
影像检查	89.07	91.72	90.38
疾病和诊断	87.52	90.91	89.18
手术	47.83	80.49	60.00
药物	89.11	99.33	93.94
检验	83.30	78.35	80.75
综合	87.29	89.18	88.22

表8 CMeEE 各个实体类别的实验结果（%）

类型	CMeEE		
	P	R	F1
疾病	77.59	77.68	77.64
症状	63.73	56.82	60.08
药物	69.08	82.03	75.00
医疗设备	58.06	72.00	64.29
医疗程序	49.77	61.08	54.85
解剖部位	62.93	73.57	67.84
检查和检验	57.60	28.57	38.20
微生物类	74.67	72.73	73.68
科室	57.14	57.14	57.14
综合	65.39	66.92	66.14

4 讨论与临床价值分析

本文模型在四个公开数据集上取得了优异性能，表明模型能为电子病历结构化及临床辅助决策场景提供切实的支持。第一，支撑细粒度医学知识图谱构建，嵌套实体在电子病历中普遍存在，传统扁平化标注方式难以刻画概念间的从属关系，本文模型能有效识别嵌套结构，如将“头颅X线片”整体识别为检查，同时将其中的“头颅”识别为解剖部位，从而为知识图谱提供更精准的语义边界和层级结构，有助于构建更符合临床认知逻辑的医学知识体系。第二，提升临床决策支持系统的智

能化水平，基于实验结果，模型对药物、疾病等高频实体的识别精度较高，能将非结构化的病历文本转化为结构化信息单元，这些高置信度的结构化数据可输入临床决策支持系统的推理引擎，以匹配并检索相关诊疗规范、药品说明或检查指南，为医生提供参考性候选列表及实时辅助提醒，从而有效减少因关键实体遗漏或语义误判导致的临床决策偏差。

但实验结果也表明模型在实际应用中仍存在明显局限。不同实体类别间的识别性能存在差异，部分低频类别和长复合型实体的识别精度仍有不足，说明模型在当前阶段尚不适合完全替代人工审核。因此，该模型更适合作为电子病历结构化流程中的前置预处理模块，而非终端决策工具。在实际部署时，需将模型输出与人工审核相结合，形成“自动提取—置信度标注—人工确认”的处理流程：对模型预测置信度较高的实体结果可直接采纳，对低置信度或低频类别的识别结果标记为待审核状态，交由专业人员复核确认。通过上述方式，模型的高效提取能力与人工的专业判断能够形成有效互补，从而保证临床信息处理的准确性和可靠性。

5 结语

针对电子病历命名实体识别面临的术语专业性强、标注数据稀缺及实体嵌套结构复杂等问题，本文提出一种融合外部知识与提示学习的电子病历命名实体识别模型。模型利用医学先验知识构建可学习领域提示向量，在编码阶段融合外部词典知识，在解码阶段使用“先边界、后类别”的两阶段识别框架。在四个公开数据集上的对比实验和消融实验证明了模型的有效性。

未来将从两方面重点开展研究：一方面，持续改进模型以提升识别性能与泛化能力，对比更多医学领域预训练模型和大语言模型，明确不同模型范式在电子病历命名实体识别任务中的适用性与互补性。此外，针对目前注意力融合外部词典无法实现词汇级别细粒度去噪的问题，将探索“引入词汇级置信度评估”和“基于对比学习的噪声过滤”两个技术路线，进一步提升模型在词典质量不完美场景下的鲁棒性。另一方面，着力推动模型在真实医学场景中的应用落地，将其与医院信息系统或临床决策支持系统深度融合，为电子病历结构化处理、医学知识图谱构建及临床辅助决策提供智能化支撑，赋能医学信息化建设，助力智慧医疗发展。

作者贡献：李阳负责实验实施与分析、论文撰写与修订；肖万里负责数据收集、清洗、预处理；赵文涵负责模型设计、可视化展示；曹子佳负责提供指导、论文审核利益证明：所有作者均声明不存在利益冲突。

参考文献

[1] 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(1):

- 15-27.
- [2] 丁建平,李卫军,刘雪洋,等.命名实体识别研究综述[J].计算机工程与科学,2024,46(7):1296-1310.
- [3] 吉旭瑞,魏德健,张俊忠,等.中文电子病历信息提取方法研究综述[J].计算机工程与科学,2024,46(2):325-337.
- [4] 陈婕卿,竹志超,张锋,等.中文电子病历命名实体识别方法研究[J].医学信息学杂志,2024,45(4):78-84.
- [5] 许思特,孙木.基于命名实体识别与 Neo4j 的中文电子病历知识图谱构建和应用[J].医学信息学杂志,2022,43(12):50-56.
- [6] 吴欢,车贺宾,王万玲,等.基于循证医学和电子病历数据的通用医学知识图谱构建[J].医学信息学杂志,2025,46(2):22-28.
- [7] 朱声荣,计虹.应用大语言模型深化临床辅助决策支持系统建设实践[J].中国数字医学,2026,21(03):24-29.
- [8] 吕婷钰,李晓瑛,张颖,等.中文医学知识大模型问答语料数据集构建研究[J].医学信息学杂志,2024,45(05):20-25.
- [9] 崔晓笛,吴冠朋,刘文强.基于 BERT-CNN-SIFRank 的智能预问诊模型研究与设计[J].中国数字医学,2025,20(8):65-71.
- [10] SAVOVA G K, MASANZ J J, OGREN P V, et al. Mayo clinical text analysis and knowledge extraction system (cTAKES): architecture, component evaluation and applications [J]. Journal of the american medical informatics association, 2010, 17 (5): 507 - 513.
- [11] Lee K J, Hwang Y S, Kim S, et al. Biomedical named entity recognition using two-phase model based on SVMs[J].Journal of biomedical informatics, 2004, 37(6): 436-447.
- [12] Jingru L, Yang S, Rui J, et al. A BiLSTM-CRF model for protected health information in Chinese[J]. Data analysis and knowledge discovery, 2020, 4(10): 124-33.
- [13] Devlin J, Chang M W, Lee K, et al. Bert: Pre-training of deep bidirectional transformers for language understanding[C]//Proceedings of the 2019 conference of the north american chapter of the association for computational linguistics: human language technologies, volume 1 (long and short papers). 2019: 4171-4186.
- [14] Liu Y, Liu M. Research on Named Entity Recognition of Traditional Chinese Medicine Text Based on RoBERTa-BiLSTM-CRF[C]//2024 9th International conference on intelligent informatics and biomedical sciences (iciibms). 2024, 9: 130-135.
- [15] Ma R, Peng M, Zhang Q, et al. Simplify the usage of lexicon in Chinese NER[C]//Proceedings of the 58th annual meeting of the association for computational linguistics. 2020: 5951-5960.
- [16] He Q, Chen G, Song W, et al. Prompt-based word-level information injection BERT for Chinese named entity recognition[J]. Applied sciences, 2023, 13(5): 3331.
- [17] Mai C C, Chen Y, Gong Z, et al. PromptCNER: A segmentation-based method for few-shot Chinese NER with prompt-tuning[J]. ACM transactions on asian and low-resource language information processing, 2025, 24(9): 1-24.
- [18] Li X, Yan H, Qiu X, et al. FLAT: Chinese NER using flat-lattice transformer[C]//Proceedings of the 58th annual meeting of the association for computational linguistics. 2020: 6836-6842.
- [19] An Y, Xia X, Chen X, et al. Chinese clinical named entity recognition via multi-head self-attention based BiLSTM-CRF[J]. Artificial intelligence in medicine, 2022, 127: 102282.
- [20] Luo Z, Che J, Ji F. Chinese medical named entity recognition based on multimodal information fusion and hybrid attention mechanism[J]. Plos one, 2025, 20(6): e0325660.
- [21] Su J, Murtadha A, Pan S, et al. Global pointer: Novel efficient span-based approach for named entity recognition[J]. Arxiv preprint arxiv:2208.03054, 2022.
- [22] Li J, Fei H, Liu J, et al. Unified named entity recognition as word-word relation classification[C]//Proceedings of the aaai conference on artificial intelligence. 2022, 36(10): 10965-10973.
- [23] Yan Y, Kang Y, Huang W, et al. Chinese medical named entity recognition utilizing entity association and gate context awareness[J]. Plos one, 2025, 20(2): e0319056.