

Cryword field named entity recognition technology based on data enhancement and low rank fine-tuning

Leer Bao Wei Zhang

Beijing Information Science and Technology University, Beijing, 102206, China

Abstract

With the growth of password security requirements, the named entity identification (Named Entity Recognition, NER) in the field of network security has been widely concerned, among which the identification of password entity has become a key difficulty. Due to the expertise and physical complexity of cryptographic data, the recognition ability of existing NER methods is limited. In addition, the scarcity of annotated data in cryptographic field makes it difficult to train NER models. To address these problems, a CRYPT-DALoRA is proposed. First, the initial dataset is built with business data as data sources, and the data augmentation method is used to construct the NER data set based on thinking chain prompt fine-tuning. Secondly, based on the thinking chain method, NER prompt word templates are designed for business entities, cryptographic entities and nested cryptographic entities respectively, which fine-tune the large model by combining the low rank adaptation fine-tuning method based on weight decomposition and differential learning rate. Finally, experiments were performed on the constructed cryptodomain NER dataset. This method achieved 76.87% F1-score, a 10.96 percentage points improvement compared to BERT-BiLSTM-CRF.

Keywords

large language model; named entity recognition; data enhancement; low rank fine-tuning; cryptography

基于数据增强和低秩微调的密码领域命名实体识别技术

包乐尔 张伟

北京信息科技大学, 中国 · 北京 102206

摘 要

随着密码安全需求的增长, 网络安全领域的命名实体识别 (Named Entity Recognition, NER) 受到广泛关注, 其中密码实体的识别成为关键难点。由于密码领域数据的专业性和实体复杂性, 导致现有NER方法的识别能力受限, 此外, 密码领域的标注数据稀缺也使得训练NER模型变得困难。针对这些问题, 提出了一种基于数据增强和低秩微调的密码实体识别方法 CRYPT-DALoRA。首先, 以业务数据等为数据源构建初始数据集, 并采用基于思维链提示微调的数据增强方法构建上下文语义一致的密码领域NER数据集。其次, 基于思维链方法分别设计了针对业务实体、密码实体和嵌套密码实体的3类密码领域NER提示词模板, 结合基于权重分解和差异学习率的低秩适应微调方法对大模型进行微调。最后在构建的密码领域NER数据集上进行了实验。本方法取得76.87%的F1-score, 相比BERT-BiLSTM-CRF有10.96个百分点的提升。

关键词

大语言模型; 命名实体识别; 数据增强; 低秩微调; 密码学

1 引言

命名实体识别 (Named Entity Recognition, NER) 是自然语言处理中的一项重要任务, 旨在从文本中识别关键实体并将其分类到预定类别。

随着互联网与数字化进程的加速, 密码领域数据的复杂性和规模大幅增长。而传统人工处理方式效率低下、难以满足需求。命名实体识别 (Named Entity Recognition, NER) 技术通过精准识别加密算法、密码设备等实体, 可提升信息检索与分析效率, 支持智能决策并推动工作流程自动

化。因此, 开发适配密码领域特性的 NER 技术已成为研究焦点。

随着大模型技术的快速发展, 如 GPT^[20] 等性能卓越的模型在通用领域展现出强大的能力。然而, 密码领域的 NER 任务仍面临多重挑战: 其一, 该领域实体类型复杂多样、长度多变且边界模糊, 实体嵌套情况多且复杂, 且包含大量专业术语与领域知识, 导致现有模型误判率高。其二, 高质量标注数据稀缺, 且隐私保护要求限制了数据规模, 影响全参数微调效果。其三, 密码安全场景对准确率与召回率的要求远超通用领域, 暴露了通用大模型在专业场景的局限性。这些问题严重制约了 NER 技术在密码领域的落地应用, 亟需针对性方法突破瓶颈。

【作者简介】包乐尔 (1998-), 男, 蒙古族, 中国内蒙古通辽人, 硕士, 从事计算机技术研究。

为此,本文提出一种基于数据增强和低秩微调 (Low-Rank Adaptation of LLMs,LoRA)^[21] 密码领域 NER 方法 CRYPT-DALoRA。该方法通过多层次数据增强缓解数据稀缺问题,在少量初始数据集的基础上生成高质量数据。并结合思维链 (Chain of thought,CoT)^[22] 技术和权重分解的低秩微调策略,提升大模型在密码领域 NER 任务中的表现。本文主要贡献如下:

- (1) 提出一种基于思维链提示微调的密码领域数据增强方法。针对密码领域的数据特点和 NER 任务需求,通过多层次多粒度的增强方法有效缓解高质量数据稀缺的问题。
- (2) 提出一种面向密码领域 NER 任务的模型微调方法。将密码领域实体分解为业务,密码,嵌套密码三类,分别设计 CoT 提示词模板进行指令调优,并与提出的基于权重分解和差异学习率的低秩微调技术结合,提升模型的识别能力。
- (3) 以密码领域真实业务数据,政策文件和教学资料等为数据源,构建 NER 初始数据集,并利用本文的数据增强方法增强获得了包含 2292 条上下文语义一致且专业性极强的密码领域 NER 标注数据集,最后在此数据集上进行的

对比实验,本方法取得 76.87% 的 F1-score,相比目前主流的 BERT-BiLSTM-CRF 模型和 ChatGLM4-Plus 模型分别有着 10.96 和 11.74 个百分点的提升。【1】

2 基于思维链提示微调的密码领域数据增强方法

密码领域语义环境复杂。传统数据增强方法,如随机替换或模板生成,常忽视实体嵌套和专业性,导致生成数据多样性不足、语义失真,同时面临上下文一致性不足及敏感数据隐私安全难以保障的问题。为此,本文提出基于思维链提示微调的密码领域数据增强方法,通过结构化推理和领域知识注入提升数据质量。【2】

如图 1 所示,本方法包含四层增强流程,具体而言:首先,实体标记层通过规则标注为后续增强奠定基础,提升模型对指定实体的识别能力。其次,数据脱敏层通过脱敏处理确保敏感数据的安全。再次,基于密码领域知识增强的实体替换层注入密码领域知识,确保替换后文本的专业性和上下文一致性。最后,基于密码领域知识增强的多粒度语义增强层在注入密码领域知识后,通过词汇替换、句子插入和段落重组等方法,从词汇至段落层面提升数据多样性。

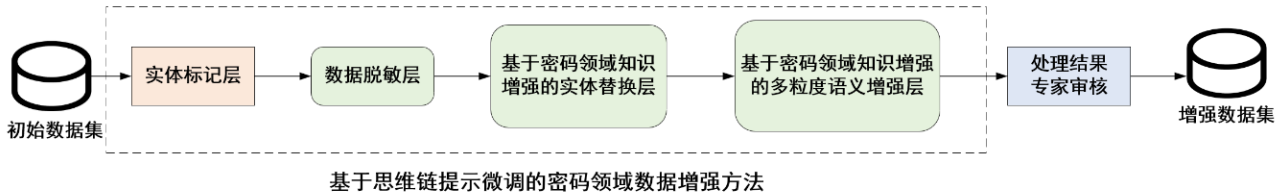


图 1 数据增强方法流程图

2.1 实体标记层

为保证数据质量,本研究以真实业务数据、政策文件和教学资料为数据源,通过专家采样筛选和预处理数据,构建初始数据集。基于对数据特征的分析,经与密码领域专家研讨,定义了八个密码领域实体类别,在表 3.1 中提供了各

类实体的类型、标签和含义。

在此基础上,设计了一个基于规则的实体标记器,通过正则表达式在初始数据集中为八类实体添加标签。算法按实体名称长度从长到短排序,以优先处理嵌套实体,避免短实体标签覆盖长实体,确保标注准确性。【3】

表 1 密码领域实体定义表

实体类型	标签	实体含义
系统	sys	在密码场景中,由硬件和人员等要素组成,用于实现特定密码功能的整体。
单位	co	涉及密码的各组织机构,如政府部门、企业、事业单位等。
机房	rm	专门用于放置密码设备、服务器等硬件设施的场所。
密码设备	dev	用于生成、存储、加解密、认证等密码操作的硬件设备。
密码算法	alg	通过数学计算实现加密、解密、认证等密码功能的一系列规则和步骤。
密码协议	prot	为实现安全通信、密钥交换、身份认证等密码相关功能,在网络环境中各参与方之间遵循的一系列协定和规则。
通信信道	ch	通信信道在密码领域是用于传输加密数据或密钥等信息的通信路径。
密码证书	cert	由权威颁发机构颁发,用于在网络环境中证明用户身份和公钥的真实性。

2.2 数据脱敏层

初始数据集来源于真实业务数据和政策文件,涉及敏感密码安全信息,直接使用存在泄露风险。为此,设置数据

脱敏层,将敏感数据转化为安全可用数据。数据分析表明,敏感信息主要集中在系统、单位、机房和通信信道实体中,因其通常包含具体名称或位置信息,故本层针对这四类实体

进行脱敏处理。

传统脱敏方法基于规则替换或删除敏感内容，导致数据重复率高、多样性差，与真实场景差异较大。为克服这一局限，在实体标记层输出的基础上，设计了数据脱敏思维链提示模板 P_{obf} ，如公式 (3.1) 所示，该模板分为角色定位 P_{c1} 、思维链任务描述 P_{cot1} 、生成示例 P_{os} 和待脱敏内容 T_1 部分。如公式 (3.2) 所示，上述内容共同构成本层模型输入内容 $Input_1$ ，交由大模型进行数据脱敏工作。

$$P_{obf} = P_{c1} + P_{cot1} + P_{os} \# \quad (3.1)$$

$$Input_1 = P_{obf} + T_1 \# \quad (3.2)$$

2.3 基于密码领域知识增强的实体替换层

为进一步提升数据多样性并保持专业性和上下文一致性，本文设计了基于密码领域知识增强的实体替换层。传统方法常采用随机替换或通用同义词替换，如将“SM4 算法”替换为“加密方法”，难以满足专业性要求，且易导致语义失真。【4】

为此，本层结合密码领域知识库和思维链设计替换策略。首先，构建了密码领域知识库，包含 50 种密码算法（如 SM4、AES、SHA1 等）、40 种密码设备（如智能密码钥匙、签名验签服务器等）、数十种密码协议和证书及其功能属性介绍，并由专家验证准确性。其次，设计了实体替换提示词模板 P_{rep} ，如公式 (3.3) 所示，该模板分为角色定位 P_{c2} 、密码领域背景知识注入 P_{k1} 、思维链任务描述 P_{cot2} 、处理结果检查 P_{e1} 、生成示例 P_{fs1} 五部分。最后拼接待实体替换的文本，即数据脱敏层的处理结果。如公式 (3.4) 所示，上述内容共同构成实体替换层的大模型输入。

$$P_{rep} = P_{c2} + P_{k1} + P_{cot2} + P_{e1} + P_{fs1} \# \quad (3.3)$$

$$Input_2 = P_{rep} + T_2 \# \quad (3.4)$$

该方法生成的替换数据在保留密码领域专业特性的同时，显著提升了数据的多样性。具体而言，通过结合密码领域知识库和思维链提示技术，替换后的数据不仅维持了实体（如密码算法、设备等）的专业术语准确性和上下文语义一致性，还通过多样化的实体替换（如将“SM4 算法”替换为“AES 算法”）扩展了数据集的覆盖范围，使其能够更全面地反映密码领域中不同场景和实体的特征。

2.4 基于密码领域知识增强的多粒度语义增强层

为进一步提升数据多样性和语义丰富性，同时保留数据的密码领域专业性，本文设计了基于知识增强的多粒度语义增强层。传统方法局限于词汇替换，如简单地调整“系统使用 SM4 加密”为“系统使用 SM4 编码”，增强数据的专业性和稳定性较低，且无法反映句子和段落层面的变化，难以满足对增强数据多样性的要求。

语义增强层基于思维链提示词模板进行工作，将增强工作细分为词汇、句子和段落三个粒度，其知识注入继承沿

用了实体替换层构建的密码领域背景知识库。基于密码领域知识增强的多粒度语义增强提示词模板 P_{aug} 具体组成如公式 (3.5) 所示，该模板主要包含角色定位 P_{c3} 、密码领域背景知识注入 P_{k2} 、思维链任务描述 P_{cot3} 、处理结果检查 P_{e2} 、生成示例 P_{fs2} ，从不同粒度增强文本语义。【5】

$$P_{aug} = P_{c3} + P_{k2} + P_{cot3} + P_{e2} + P_{fs2} \# \quad (3.5)$$

本层的输入文本内容为提示词模板拼接上实体替换层的输出结果。上述各部分共同构成公式 (3.6) 中的实体替换层的大模型输入，交由大模型进行语义增强工作。

具体而言，本层工作在词汇层面，采用两种策略：一是根据文本语义，随机插入密码领域专业词汇（如“MD5+Salt 算法”），增强专业性并提升生成质量。二是对非关键实体进行同义词替换或词汇删除，增加表达多样性，再由大模型修正至通顺。在句子层面，结合文本语义和知识库的 ICL 能力，插入上下文一致的密码领域句子（如“系统通过 MD5+Salt 算法保护关键数据”），并通过随机拆分与组合挖掘语义潜力，后由大模型优化至流畅。在段落层面，运用随机拆分与合成增强语义多样性，扩大生成数据的规模与多样性。

$$Input_3 = P_{aug} + T_3 \# \quad (3.6)$$

相较于传统的基于规则的语义增强方法，本文提出的基于思维链提示微调的密码领域数据增强方法，不仅有效扩充数据规模，还显著提升增强数据的质量与专业性。增强后的文本在语义一致性、专业性及多样性方面均有明显改善。

3 面向密码领域命名实体识别任务的大模型微调训练方法

3.1 基于思维链的密码实体识别指令微调

本节定义密码领域 NER 任务并对目标实体进行了分类。密码领域 NER 任务以待识别的文本为输入，输出 JSON 格式的实体识别结果（如 { "密码算法": "SM4", "密码设备": "堡垒机" } ），用于识别与密码安全相关的实体（如密码算法、密码设备等）。结合对密码领域 NER 的数据和任务的分析，将实体分为三类：业务实体、密码实体和嵌套密码实体。分别设计提示词模板。

首先是业务实体，像系统、单位和机房这类在通用文本中同样较为常见的实体，模型虽在预训练时获得了一定的识别能力，但这些实体在密码场景下有着一定的独特性，其文本中的实体含义和上下文关系与通用场景存在差异。例如，在密码领域，“系统”指实现特定密码功能的整体，这与通用语境下的“系统”含义不同。所以，针对密码场景对模型进行针对性优化，能够确保模型准确识别这些实体，避免因实体领域差异导致的误判。

然后是密码实体，涵盖加密算法、密码协议和密码设备这三类密码领域专业实体。识别这些实体需要进行密码领域知识注入的增强，因此在微调过程中，在模板中加入了密

码实体识别的学习示例，同时，从预先构建好的密码领域背景知识库中选取注入了详细的知识，以强化模型对专业实体的语义理解和特征提取能力。

最后是识别难度最高的嵌套密码实体，这例如“本地 IPsec VPN 通过基于 RSA2048 算法和 SHA-256 算法的数字证书进行验证，数字证书为自签名方式”，涉及多个密码算法、密码证书等的嵌套关系，需要模型对密码领域知识的深度理解能力。将这类实体单独作为一部分进行处理，设计专用的 CoT 提示词模板，有助于模型逐步学习和掌握处理复杂嵌套结构密码实体的能力，解决在处理此类实体时识别效果不佳的问题。

3.2 基于权重分解和差异学习率的低秩模型方法

密码领域 NER 任务因实体嵌套和专业术语复杂性较高，

传统 LoRA 微调在特征适配和效率上需进一步提升。为此，如图 2 所示，本方法在 DoRA^[23] 的基础上对微调方法进行了优化，通过引入差异学习率的方法进一步提升微调效果和训练效率。DoRA 将预训练权重分解为幅度和方向两部分，仅更新方向的低秩矩阵 A 和 B，以降低推理开销并提高性能。然而，DoRA 对 A 和 B 采用相同学习率，忽略其对输入输出特征的不同贡献，导致特征学习和训练效率不足。为解决此问题，依据对方向权重低秩适应得到的矩阵 A 和 B 的深入理解，结合 LoRA+^[24] 在原 LoRA 上设置差异倍数学习率实现性能提升的思路，为方向权重的低秩适应矩阵 A 和 B 设置差异学习率，设置 B 的学习率为 A 的 8 倍，进一步加速特征学习的效率。

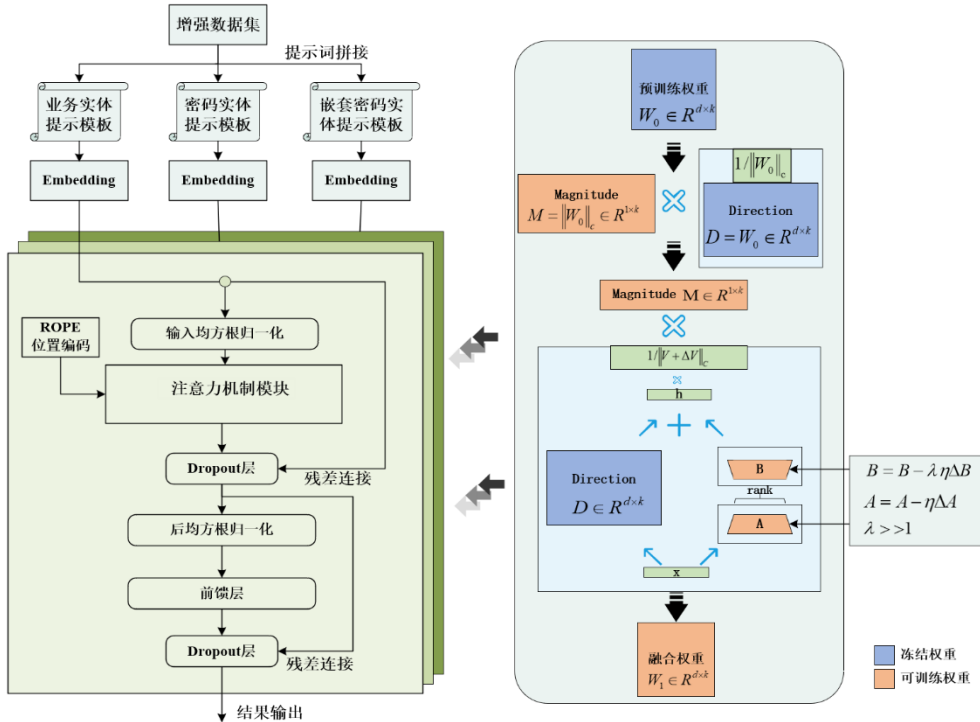


图 2 基于权重分解和差异学习率的低秩微调方法架构图

方法分为以下三个部分，首先是权重分解部分，如公式 (3.7) 所示，将预训练权重 w_0 分解为幅度权重 M 和方向权重 D 两部分。

$$W_0 = M \frac{D}{\|D\|_c} = \|W_0\|_c \frac{W_0}{\|W_0\|_c} \# (3.7)$$

然后是方向矩阵的低秩化，在方向向量的更新中，使用低秩适应的策略进行方向向量的更新，融合权重 w_1 如公式 (3.8) 所示，将方向向量适应为 A 和 B 为两个低秩矩阵， $\Delta D = BA$ 。r 为 LoRA 方法中的超参数 rank，指的是对模型参数进行低秩分解时所选择使用的秩， $r \ll \min(d, k)$ 。

$$W_1 = M \frac{D + \Delta D}{\|D + \Delta D\|_c} = M \frac{D + BA}{\|D + BA\|_c} \# (3.8)$$

最后引入差异学习率机制，针对方向权重低秩适应产生的 A 和 B 两个矩阵， ΔA 和 ΔB 为矩阵 A 和矩阵 B 的梯度。 η 为基础学习率，为两个矩阵设置不同的学习率 η_a 和 η_b ，学习率比例定义为 $\eta_B = \lambda \eta_A$ ，其中， $\eta_A = \eta$ 而 $\eta_B = \lambda \eta$ 。然后，根据对实际微调结果的分析，并结合 LoRA+ 方法中的结论，设置了 $\lambda = 8$ 。其公式如下：

$$A = A - \eta_A \Delta A \# (3.9)$$

$$B = B - \eta_B \Delta B \# (3.10)$$

$$\eta_B = \lambda \eta_A, \lambda = 8 \# (3.11)$$

通过权重分解和差异学习率的结合，有效提升了模型进行特征学习的效率和效果，在数据集规模较小的情况下尤为明显。同时，相较于全参微调，本方法大幅减少了微调

过程中需要训练的参数数量，有效降低时间开销和对算力的要求。

4 实验

4.1 数据集

实验使用的初始数据集由真实业务数据，密码领域政策法规文件，密码领域教学材料等组成，真实业务数据来自于密评机构，包括密评报告，建设方案等。在这些文件中选取与密码 NER 相关度高且覆盖任务场景的内容应用本文提出的数据增强方法获得包含 2292 条高质量数据的密码领域 NER 标注数据集，确保了数据的多样性和覆盖性。最后将数据集按照 8: 1: 1 的比例划分为训练集、验证集和测试集，用于实验。

4.2 评价指标

Precision: 精确率衡量的是在所有被预测为正类的样本中，实际为正类的比例。计算方法如下公式 (3.12) 所示。

$$precision = \frac{Num_{TP}}{Num_{TP} + Num_{FP}} \# (3.12)$$

Recall: 召回率衡量的是在所有实际为正类的样本中，被模型正确预测为正类的比例。计算方法如下公式 (3.13) 所示。

$$recall = \frac{Num_{TP}}{Num_{TP} + Num_{FN}} \# (3.13)$$

F1-score: F1 分数是精确率和召回率的调和平均数。计算方法如公式 (3.14) 所示。

$$F1 - score = \frac{precision \times recall \times 2}{precision + recall} \# (3.14)$$

4.3 实验环境

实验环境为：120G 内存，GPU 为 NVIDIA GeForce RTX 4090，CPU 为 Xeon(R) Gold 6430，显存为 24GB，CUDA 12.4，运行环境为 python3.10。参数如表 3 所示。

表 3 实验参数设置表

参数名称	参数值
Learning Rate	1e-5
Epoch	3
Batchsize	16
LoRA-rank	16
LoRA-alpha	32

4.4 对比实验

为了全面评估所提出方法的性能和效果，设计了对比实验，选取以下模型进行对比分析：

Qwen: 是由阿里巴巴开发的一系列开源模型，效果优异，各项能力均在国内外模型前列，实验选择 Qwen 系列中的 Max 和 Plus 两款模型进行对比。

ChatGPT: 是由 OpenAI 团队开发的一系列高水平模型，

是目前业界主流的大模型之一。实验选择 GPT3.5-turbo 和 GPT4.0 两款模型进行对比。

ChatGLM4: 是由智谱华章公司开发的开源大模型系列，中文处理能力极为突出，实验选择该系列中的 Plus、0520、Air 和 9B 四款进行对比。

BERT-BiLSTM-CRF: 模型基于 BERT，结合了 BiLSTM 和 CRF 的能力和优势，在 NER 任务中取得了优秀的效果，是 NER 任务的主流模型。

对比实验结果如表 4 对比实验成绩表所示，各数据保留两位小数。

表 4 对比实验成绩表

模型	精准率	召回率	F1 分数
Qwen-Plus	41.50	42.28	41.87
ChatGPT3.5-Turbo	46.79	39.19	42.65
ChatGLM4-Air	47.04	46.49	46.76
ChatGPT4.0	58.33	56.47	57.38
ChatGLM4-0520	57.53	58.69	58.10
Qwen-Max	61.73	64.21	62.95
ChatGLM4-Plus	64.37	65.91	65.13
BERT-BiLSTM-CRF	64.16	67.75	65.91
CRYPT-DALoRA	77.97	75.81	76.87

CRYPT-DALoRA 表现最为优异，F1 分数为 76.87%。这主要由于数据增强和模型微调策略的有效结合。本文的数据增强方法通过多层次粒度的处理，显著提升了数据集的规模和质量，为模型训练提供了丰富且高质量的数据。同时，提出的微调方法进一步提升了模型对复杂专业领域实体的识别能力。

可以看出与 CRYPT-DALoRA 相比，其他主流模型的表现仍有提升空间。ChatGLM4-Plus、Qwen-Max、ChatGLM4-0520 和 ChatGPT4.0 的 F1 分数分别为 65.13%、62.95%、58.10% 和 57.38%，均显著低于 CRYPT-DALoRA。这些模型由于基于通用数据训练，缺乏针对密码领域有效的数据增强，导致对密码领域实体的理解和识别能力不足，未能充分挖掘领域内的关键信息。另一方面，CRYPT-DALoRA 方法相比 BERT-BiLSTM-CRF，在 F1 分数上有 10.96 个百分点的优势，这主要是由于本方法基于任务分解的指令微调和基于权重分解和差异学习率的低秩模型微调方法。前者使模型能够更加不受干扰、精准地学习和识别各类实体的特征。后者使模型在优化过程中能更高效地学习特征，提高在较小数据集上的微调效果。此外，ChatGLM4-Air 和 Qwen-Plus 由于模型规模和推理能力有限，导致精准率和召回率较低。

4.5 消融实验

消融实验通过对比不同配置模型的性能，来探究 CRYPT-DALoRA 各组件的有效性及其相互作用。从四个配置实验的结果入手，深入分析各组件对模型性能的影响。消融实验的结果如表 3.5 所示，各数据保留两位小数。

(1) LLM+CoT: 该配置使用 ChatGLM4-9B 作为基础模型, 仅通过密码领域 NER 思维链提示词模板进行增强。

(2) LLM+LoRA: 该配置同样使用 ChatGLM4-9B 作为基础模型, 仅使用密码领域 NER 任务增强数据集通过 LoRA 方法对模型进行训练。

(3) LLM+CoT+LoRA: 该配置同时使用密码领域 NER 提示词模板和 LoRA 微调进行增强。

(4) CRYPT-DALoRA: 该配置使用 ChatGLM4-9B 作为基础模型使用本文提出的完整方法进行增强。

表 5 消融实验结果

模型	精准率	召回率	F1 分数
LLM+CoT	21.42	17.46	19.24
LLM+LoRA	76.27	72.86	74.53
LLM+CoT+LoRA	69.70	64.61	67.06
CRYPT-DALoRA	77.97	75.81	76.87

LLM+CoT 配置下 F1 分数 19.24%。表明仅使用提示微调时模型表现最差。提示微调旨在通过设计特定提示词增强模型, 但密码领域实体复杂, 专业性高, 单纯的提示微调无法充分挖掘数据特征, 而且该模型规模过小, 无法有效地利用复杂的基于思维链的密码领域 NER 提示词模板提升任务表现, 导致该配置在密码实体识别任务中表现不佳。

LLM+LoRA 配置下 F1 分数 74.53%, LoRA 通过低秩矩阵对模型权重进行调整, 在处理密码领域数据时, 可以使模型得到明显增强, 对密码领域各类实体识别有更好的表现, 但没有充分利用到大模型的上下文学习能力, 导致其性能仍有提升空间。

LLM+CoT+LoRA 的 F1 分数 67.06%。在应用 LoRA 微调的基础上结合基于思维链的密码领域 NER 提示词模板后, 模型表现没有提升反而下降, 是因为 ChatGLM4-9B 模型规模较小, 基于思维链、专业领域知识注入和 few-shot 的提示词模板对其而言过于复杂, 无法理解并利用提示词模板, 反而被其中的内容混淆导致表现降低。

使用本文完整方法时 F1 分数 76.87%, 由于本方法不仅包含针对密码领域数据特点的数据增强策略, 为模型提供丰富且高质量的数据, 还将密码领域 NER 任务分解, 并为

各类型实体分别设计提示词模板, 使模型能更好地对不同类型实体进行学习和处理, 结合基于权重分解和差异学习率的低秩模型微调方法改善特征学习的效果和效率。各方法相互配合, 有效提升了模型在密码领域实体识别任务中的成绩。

5 总结

本文针对密码领域 NER 任务中数据稀缺、专业性强、实体复杂且嵌套等挑战, 提出了一种基于数据增强和低秩微调的 CRYPT-DALoRA 方法, 显著提升了大模型在该领域的识别性能。通过构建密码领域 NER 初始数据集、设计多层次多粒度数据增强策略及优化的大模型微调方法, 本文为密码领域 NER 任务提供了高效可靠的解决方案。

首先, 以真实业务数据、政策文件和教学资料为源, 通过专家筛选与规则标记, 梳理出八类关键实体, 并依此构建初始 NER 数据集。其次, 提出基于思维链提示的四层数据增强方法, 生成包含 2292 条高质量标注数据的数据集, 有效缓解数据稀缺问题并有效提升数据集的专业性与多样性。再次, 将实体分为业务实体、密码实体和嵌套密码实体三类, 设计针对性的思维链提示模板, 结合 ChatGLM4-9B 模型, 通过少样本学习与领域知识注入优化识别能力。最后, 引入基于权重分解和差异学习率的低秩微调策略显著提升模型效率与鲁棒性。

实验结果验证了方法在密码领域 NER 任务中的优越性。为密码领域 NER 任务提供了有力支持。

参考文献

- [1] Floridi L, Chiriatti M. GPT-3: Its nature, scope, limits, and consequences[J]. Minds and Machines, 2020, 30: 681-694.
- [2] Hu E J, Shen Y, Wallis P, et al. Lora: Low-rank adaptation of large language models[J]. arXiv preprint arXiv:2106.09685, 2021.
- [3] Wei J, Wang X, Schuurmans D, et al. Chain-of-thought prompting elicits reasoning in large language models[J]. Advances in neural information processing systems, 2022, 35: 24824-24837.
- [4] Liu S Y, Wang C Y, Yin H, et al. Dora: Weight-decomposed low-rank adaptation[J]. arXiv preprint arXiv:2402.09353, 2024.
- [5] Hayou S, Ghosh N, Yu B. Lora+: Efficient low rank adaptation of large models[J]. arXiv preprint arXiv:2402.12354, 2024.