

基于大语言模型的社交媒体中文命名实体识别方法

蒋大锐, 贺敏伟, 徐胜超

(广州华商学院 人工智能学院, 广州 511300)

摘要: 社交媒体中的中文文本因其类别模糊等特性, 导致文本特征提取困难, 进而影响了命名实体识别的准确性; 为此提出基于大语言模型的社交媒体中文命名实体识别方法; 该方法从社交媒体平台采集大量中文文本数据, 并对其进行语义增强处理, 随后设计文本信息编码器, 对增强后的文本进行相对位置编码; 在此基础上, 利用大语言模型深度挖掘文本信息特征, 构建特征向量表示; 同时, 引入注意力机制, 分析文本上下文的语义特征; 通过标签解码过程, 输出最优的文本序列, 从而构建中文命名实体识别模型; 在损失函数的指导下, 对模型参数进行优化, 实现对命名实体的精确识别; 实验结果表明, 该方法在实际应用中表现出 0.03% 的误识率, 具备较高的识别精度和良好的应用效果。

关键词: 大语言模型; 社交媒体; 中文命名实体; 实体识别; 文本编码; 识别模型; 标签解码

Chinese-Named Entity Recognition Method for Social Media Based on Large Language Model

JIANG Darui, HE Minwei, XU Shengchao

(School of Artificial Intelligence, Guangzhou Huashang College, Guangzhou 511300, China)

Abstract: It is difficult for Chinese text in social media to extract text features because of its fuzzy category, which affects the accuracy of named entity recognition. Therefore, a large language model based Chinese-named entity recognition method for social media is proposed. This method collects a large number of Chinese text data from social media platforms, enhances the semantic processing of the data, and then designs a text encoder to encode the relative position of the enhanced text. On this basis, the large language model is used to deeply explore the text information features and construct the feature vector representation, and the attention mechanism is introduced to analyze the semantic features of the text context. Through the label decoding process, the optimal text sequence is output to construct a Chinese-named entity recognition model. Under the guidance of the loss function, the model parameters are optimized to realize accurate identification of named entities. Experimental results show that in practical applications, this method has the error rate of 0.03%, and has high recognition accuracy and good application effect.

Keywords: big language model; social media; Chinese-named entities; entity recognition; text encoding; identification model; tag decoding

0 引言

命名实体识别作为在自然语言处理领域中极为常见且关键的部分, 其能够通过计算机去识别出现有文本中的具有现实意义的实体, 如籍贯、机构名、国籍等。这

些实体能够在文本信息的深入理解、特征挖掘等方面起到非常重要的作用, 也在语义分析、自动问答、信息抽取等领域扮演着极为重要的角色。命名实体识别的准确度对上层应用的效果好坏有着极大的影响。

尤其在处理社交媒体文本时, 命名实体识别面临着

收稿日期:2024-12-20; 修回日期:2025-02-17。

基金项目:国家自然科学基金面上项目(61972444);广州华商学院校内导师制科研项目(2024HSDS27)。

作者简介:蒋大锐(1986-),男,博士研究生,副教授。

贺敏伟(1963-),男,博士,教授,硕士生导师。

徐胜超(1980-),男,硕士,副教授。

引用格式:蒋大锐, 贺敏伟, 徐胜超. 基于大语言模型的社交媒体中文命名实体识别方法[J]. 计算机测量与控制, 2025, 33(6): 247-253.

更为严峻的挑战。社交媒体文本的语言风格多样, 构词灵活多变, 且往往存在类别模糊的问题。具体而言, 社交媒体中的中文文本常常使用口语化、网络化的表达方式, 这些表达方式往往不符合传统的语法规则, 导致传统的命名实体识别方法难以准确识别。此外, 社交媒体文本中的命名实体往往缺乏明确的边界和固定的格式, 使得识别过程更加复杂。类别模糊的问题还体现在社交媒体文本中命名实体的多样性上, 同一类实体可能以多种不同的形式出现, 而不同的实体又可能具有相似的表现形式, 这使得识别系统难以准确区分。

在该研究背景下, 不少研究学者针对实体识别展开了研究, 并已经得出一定研究成果。如文献 [1] 首先对中文文本进行词性标注, 通过深度学习编码并嵌入字符, 提取关键特征, 捕捉文本信息。利用条件随机场 (CRF, conditional random field) 标注序列信息, 解码特征后构建嵌套实体识别框架用于识别嵌套命名实体, 并进行后处理, 以提升命名实体识别效果。该方法应用的序列标记依赖于分词和词性标注的准确性, 若词性标注出现错误, 将会直接影响到后续识别结果的准确性。文献 [2] 现根据中文命名实体的特点, 设计触发器并构建触发器库, 将现有的中文文本进行归一化处理, 然后利用深度学习框架, 构建匹配网络, 提取出文本特征, 最后将其与设计的触发器进行匹配计算, 筛选出相应的实体, 并对识别结果进行校验。该方法设计的触发器有着质量和数量的限制, 若触发器库不完整, 将会导致识别结果不够准确。文献 [3] 利用预训练的词向量和字符向量捕捉词汇的语义信息, 并提取汉字的字形特征、语音特征, 然后对文本信息进行分析, 深入提取文本的上下文特征, 将上述多个特征通过加权计算进行融合, 建立特征向量, 最后在注意力机制的作用下, 对向量中的多个特征进行熵值计算, 筛选出对命名实体识别有影响的特征, 在双仿射变换的基础上, 捕捉命名实体与边界之间的关系, 从而实现对命名实体的准确识别。该方法采用了多特征融合方法, 需要提取多种特征信息, 这增加了特征提取的复杂性, 导致其识别结果不够准确。文献 [4] 构建一个包含丰富词汇信息的 Lexicon, 然后设计指针网络, 利用该网络接收文本的字符和词嵌入表示, 引入循环神经网络, 筛选出关键特征, 通过计算特征熵值, 对特征进行深入挖掘, 生成特征序列, 然后计算特征序列与 Lexicon 之间的相似度, 识别出相应的命名实体。但该方法应用的指针网络较为复杂, 对其进行构建时需要花费大量的时间, 导致方法的识别时间较长。

在以往研究的基础上, 针对存在的不足, 本文设计了一种基于大语言模型的社交媒体中文命名实体识别方法。本研究通过引入大语言模型, 深化对命名实体识别

机制的理解, 促进相关领域的发展。同时, 该方法在社交媒体、公共信息监测等多个领域应用广泛, 以为数据挖掘、智能推荐、舆情监测等任务提供有力支持。本文最后通过具体的实验数据与性能分析验证了该方法的正确性与高效性。

1 社交媒体中文命名实体识别方法设计

1.1 社交媒体中文文本编码

社交媒体中中文文本的构词较为灵活, 其类别至今的界限较为模糊。为深入理解社交媒体中文文本的深层含义^[5], 需要提取出文本中的“关键词”, 即中文实体, 以便更好地理解文本信息, 分析文本的语义结构, 进一步提高文本识别的准确性^[6]。为此, 需要收集大量的社交媒体中文文本。在收集时, 主要从社交媒体平台直接获取用户发布的文本内容, 或者在图书馆和搜索引擎中获取关键词^[7], 由此, 得到海量的中文文本信息。

采集的文本数据数量庞大, 但其存在文本稀疏的情况, 简单而言, 就是其存在的命名实体数量较少, 导致很难对其进行识别和筛选^[8]。为充分挖掘文本信息的内在联系, 需要对其进行语义增强处理, 其增强过程如下所示:

$$H = A - t \quad (1)$$

$$g = \sigma(p * H + p * w + b) \quad (2)$$

式中, H 为文本信息的相对距离信息, A 为文本信息的方向信息, t 为文本信息的相似参数, g 为文本信息语义增强处理结果, σ 为文本信息的加权参数, p 为文本信息的注意力得分, w 为文本信息的加权值, b 为文本信息的输出偏置项。

采用公式 (1) 和公式 (2) 采用的语义增强技术, 对文本信息进行了深度处理, 有效去除了多余文本量对命名实体识别的干扰, 以“杭州樱花广场”这一命名实体为例, 原始文本中包含大量关于该地点的冗余描述或非关键信息, 如“在风景秀丽的杭州市中心, 有一个广为人知的景点——杭州樱花广场, 它以其独特的樱花景观和丰富的文化活动吸引了众多游客前来观赏。”这样的描述虽然详尽, 却包含了大量不必要的噪声。通过运用语义增强技术, 能够将原始文本精炼为更加聚焦于命名实体的形式, 如“杭州樱花广场吸引了众多游客。”这样的处理不仅去除了冗余的描述, 还使得命名实体“杭州樱花广场”在文本中更加突出, 从而减少了噪声对识别结果的干扰。此外, 语义增强技术还能够进一步挖掘文本中的隐含信息, 如“樱花广场”与“杭州市中心”之间的地理位置关系, 以及“樱花景观”和“文化活动”等特征属性, 这些信息对于深入理解命名实体的含义和背景具有重要价值。

同时, 利用这一语义增强后的文本信息^[9], 设计

了一种高效的文本信息编码器, 该编码器能够更准确地捕捉和处理文本中的关键语义特征, 为后续构建高精度的文本识别模型奠定基础。其设计的编码器具体结构如图 1 所示, 该编码器通过深度学习算法对语义增强后的文本进行高效编码, 提取出对命名实体识别的特征信息。通过这样的处理流程, 优化整个文本处理和分析的过程。

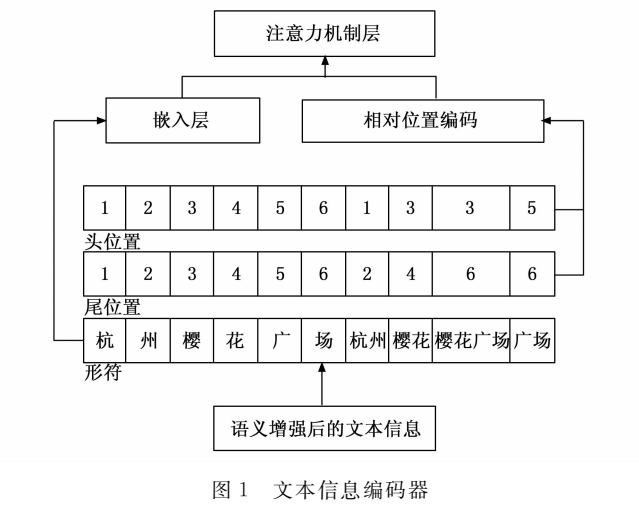


图 1 文本信息编码器

在该编码器中, 同样以“杭州樱花广场”为例, 先对该句子中的单字进行划分。同时, 将其划分为“杭州”“樱花”“樱花广场”“广场”4 个具有实际意义的实体。在这个过程中, 针对“樱花广场”是否需要划分, 往往存在争论, 对其需要是否划分为两个实体, 需要根据实际情况进行决定, 从而才能更好地理解句子的深层含义^[10]。根据确定的实体, 再对该文本进行相对位置编码。在编码时, 根据图 1 所示的形符, 对文本的头位置和尾位置分别进行相对位置编码, 由此得到唯一且准确的编码结果^[11], 为后续构建命名实体识别模型奠定基础。

通过完整的文本编码流程, 包括文本收集、语义增强处理和文本信息编码器设计, 将中文文本转化为计算机可理解的编码形式, 为后续命名实体识别模型的构建提供了坚实的基础。

1.2 基于大语言模型的中文命名实体识别模型构建

将上述中文文本编码结果作为识别模型的输入序列, 利用大语言模型, 构建中文命名实体识别模型^[12]。大语言模型作为一种基于深度学习构建的能够生成自然语言模型的强大模型, 通过能够利用神经网络结构, 处理和生成文本数据^[13], 并对复杂语言进行分析和处理。将该语言模型应用于构建中文命名实体识别模型时, 它能够有效地捕捉中文文本的语言特征, 并准确推断出文本的实体类别^[14]。该大语言模型用于中文命名实体识别时网络结构如图 2 所示。

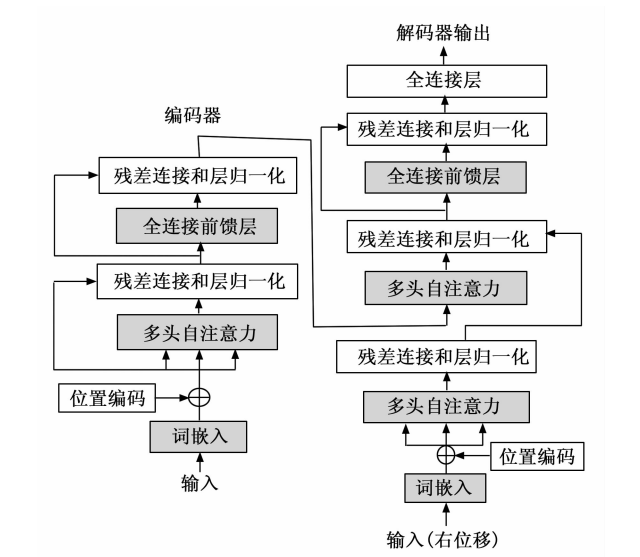


图 2 大语言模型的网络结构图

利用大语言模型, 提取出社交媒体中文文本信息的特征^[15]。其特征提取过程如下所示:

$$S = \frac{e}{k} * (\sum_{i=1}^m Q_i + \sum_{j=1}^n P_j) \tag{3}$$

式中, S 表示基于大语言模型提取的社交媒体中文文本信息特征, e 表示文本的词频概率, k 表示大语言模型的注意力参数, Q_i 表示利用大语言模型对文本数据进行映射的第 i 个子空间, P_j 表示社交媒体中文文本信息的键矩阵, m 表示映射空间的数量, n 表示文本数据的键矩阵数量。

根据提取的特征, 构建对应的特征向量, 并对特征向量进行表示, 建立相应的实体识别模型^[16], 其具体如图 3 所示。

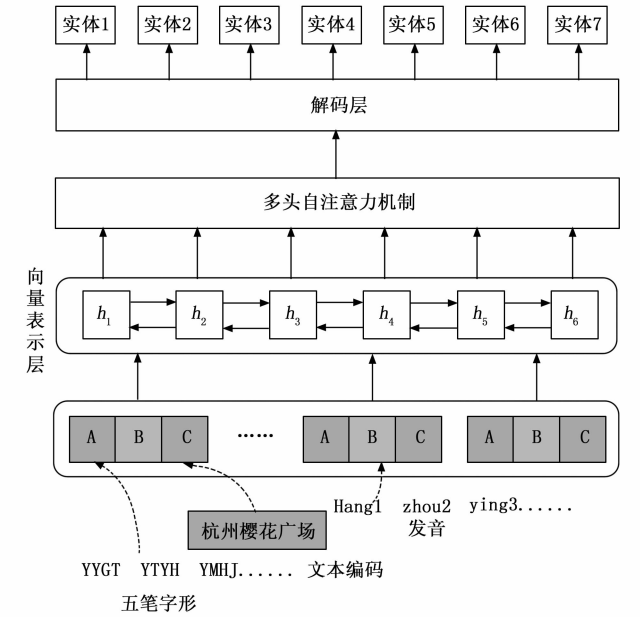


图 3 基于大模型的中文命名实体识别模型的构建过程

如图 3 所示,在该复杂的模型构建过程中,首先将图 2 中经过设计的编码算法处理后的结果作为输入变量。这一编码结果蕴含了丰富的文本信息,在模型中,对输入文本的字音和字形进行深入分析。字音分析有助于理解词汇的发音特征,而字形分析则揭示了汉字的构造和书写规律。这些分析对于捕捉中文文本的差别至关重要。接下来,结合公式(3)利用大型语言模型对输入文本进行特征提取^[17]。这些大型语言模型经过大量数据的训练,能够准确捕捉文本中的语义信息,并生成相应的特征向量。然而,直接从大型语言模型中提取的特征可能包含冗余和噪声。因此,进行了一系列特征筛选步骤,以确保只保留对中文文本识别有用的特征。这一步骤通常涉及特征选择算法和机器学习技术的综合运用。

在构建了中文文本特征向量表示后,利用多头自注意力机制^[18]来提取文本序列的上下文特征。多头自注意力机制能够同时关注文本中的多个位置,并根据不同位置的相对重要性进行加权求和,从而更准确地捕捉文本中的上下文信息。最后,利用解码层对提取的特征进行解码。解码层通常采用循环神经网络(RNN, recurrent neural network)或长短期记忆网络(LSTM, long short-term memory)等结构,能够根据提取的特征向量生成相应的解码序列^[19]。在解码过程中,保持解码序列与输入文本的一致性,以确保识别结果的准确性。通过这一系列步骤,最终实现了对中文命名实体的识别。

在上述过程中,标签解码是获得准确识别结果的关键^[20]。其标签解码的具体过程如下所示:

$$T = \frac{e}{\sum_{s=1}^a e} \quad (4)$$

$$Y = \operatorname{argmax}[T \times S(q, y)] \quad (5)$$

式中, e 表示文本的词频概率, T 表示中文文本的特征转移矩阵, a 表示迭代次数, Y 表示标签解码的结果, $S(q, y)$ 表示各标签序列的条件概率分布, y 表示全局最优标签序列, q 表示所有可能的标签序列。基于该公式,完成对标签序列的解码,由此输出基础的中文命名实体识别结果。

通过提取文本特征、构建特征向量、引入多头自注意力机制以及标签解码等步骤,构建了一个高效、准确的中文命名实体识别模型。这一模型不仅能够有效捕捉中文文本的语言特征,还能准确识别出文本中的命名实体。

1.3 中文命名实体识别结果输出

为保证构建模型输出识别结果的准确性^[21],利用损失函数,对识别结果进行后处理,优化其性能,由此

输出更为准确的识别结果^[22]。由此,输出的识别结果如下所示:

$$Y^* = \exp(Y + \kappa L) \quad (6)$$

式中, Y^* 为基于识别模型输出的实体识别结果, Y 为标签解码的结果, L 为引入的损失函数, κ 为模型参数调整系数。

基于上述公式,将识别到的中文命名实体输出,从而对舆情监测、社交媒体数据分析等多个领域提供可靠的数据基础^[23-24],并促进相关领域的发展和进步。

通过引入损失函数对识别结果进行后处理,进一步优化了模型的性能,提高了识别结果的准确性。这一步骤的工作不仅验证了模型的实用性和有效性,也为相关领域的发展和进步提供了可靠的数据支持。

2 实验测试与性能分析

2.1 实验数据集

本次实验数据集采用多个公开数据集,分别来自多个权威机构,如 MSRA、《人民日报》等,其数据集具体如表 1 所示。

表 1 实验数据集

数据集	种类	训练集	验证集	测试集
MSRA Data	句子数量	13 500	2700	2 700
	字数	738 000	145 000	15 000
	实体数量	1 890	390	450
Weibo NER Data	句子数量	5 200	510	480
	字数	1 500 000	139 000	200 000
	实体数量	150 000	15 000	16 300
People's Daily NER	句子数量	157 000	6 300	6 900
	字数	5 000 000	48 000	56 000
	实体数量	131 000	69 500	8 800
Resume Data	句子数量	254 000	12 300	17 000
	字数	45 000	12 000	13 000
	实体数量	168 000	2 630	5 450

如表 1 所示, MSRA Data 作为研究所提供的公开数据集,具有极高的权威性,其数据量庞大,实体类别相对较少,主要为机构、人名及地区名 3 大类; Weibo NER Data 是线上网站整理的公开数据集,其数据量和实体类别均比较丰富; People's Daily NER 其来源比较权威,主要来自于国内具有权威性的媒体报道,其包含人名和地名等实体; Resume Data 来自于公开的简历数据,其实体类别相对于上述 3 个数据集较多,包括国籍、学校、专业及种族等。

2.2 实验参数设置

实验选用 Intel i9-9900K 为 CPU,其操作系统为 Linux,深度学习框架为 TensorFlow,依赖库为 Transformers。由此构建相应的实验环境。在该实验环境中,设定大语言模型参数,具体模型参数如表 2 所示。

表 2 大语言模型参数

序号	模型参数	参数设置
1	Transformer 层数	12 层
2	隐藏层数	768
3	多头注意力头数	15
4	初始学习率	0.05
5	训练轮数	50
6	Batch 大小	64
7	梯度裁剪系数	5.0
8	权重衰减	0.01

如表 2 所示,除了大语言模型参数外,还需要设定方法设计过程中的部分参数,具体如下:文本信息的加权参数 σ 为 2.45,注意力参数 k 为 3.75,最优标签序列参数 γ 为 4.15,模型参数调整系数 κ 为 2.20。

2.3 实验过程

利用本文设计方法对表 1 所示的数据集进行语义增强,再利用设计的编码器对增强后的数据集进行编码处理。在表 1 设定的模型参数下,先提取编码后的数据集文本特征,建立相应的特征向量。结合文本的字形、音形等特征,构建相应的实体识别模型。引入损失函数,对识别模型进行优化,确保其输出更为准确的识别结果。损失函数如图 4 所示。

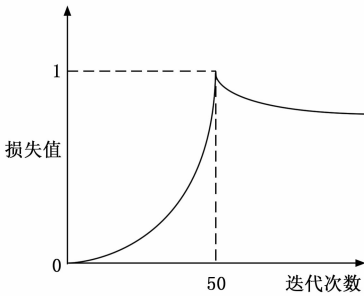


图 4 损失函数

如图 4 所示,利用该损失函数,优化识别模型中的多个参数,由此保证识别结果的准确性。

2.4 对比实验

为保证实验结果的真实性,在对本文设计方法进行验证时,设计对比实验。选择的对比组如下:文献 [1] 提出的基于序列标记的中文嵌套命名实体识别方法、文献 [2] 提出的基于触发器和匹配网络的中文命名实体识别与文献 [3] 提出的基于多特征融合和仿射的中文命名实体识别方法,3 种方法分别为对比组 1、对比组 2 和对比组 3。在对比 4 种方法的识别性能时,以误识率和 F_1 值为评价指标,对比 4 种方法的应用效果。

2.4.1 误识率对比分析

将误识率作为评价指标,对比这 4 种方法在实际应用中的效果。实验过程中,以 MSRA Data 数据集为实验对象,对其进行多次命名实体的识别,统计其识别结

果,其具体如图 5 所示。

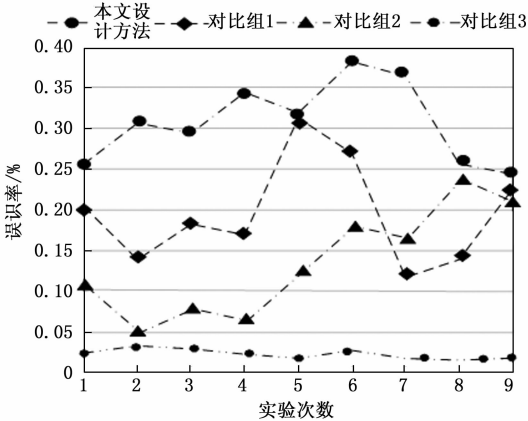


图 5 4 种方法的误识率对比分析

如图 5 所示,本文设计的方法在误识率方面表现出色,误识率均保持在 0.05% 以下,且平均误识率仅为 0.03%。这一显著优势主要得益于以下几点:首先,本文针对社交媒体中文文本的特性进行了深度优化,如考虑文本的灵活性、类别的模糊性等,使得模型更能适应社交媒体环境下的文本识别需求;其次,本文采用了基于大语言模型的算法进行命名实体识别,该技术能够更精确地捕获文本中的实体信息;再者,本文使用了高质量的数据集进行训练,确保了模型的泛化能力和识别精度;最后,采用了合理的评估方法,确保了实验结果的客观性和可靠性。相比之下,对比组的 3 种方法在误识率方面均表现不佳,均高于本文设计的方法。这可能是因为这些方法未能充分考虑到社交媒体中文文本的特性,或者采用了相对落后的技术或数据集,导致识别精度较低。

2.4.2 F_1 值对比分析

为进一步验证本文设计方法的性能,以 F_1 值为评价指标,具体的 F_1 计算公式如式 (7) 所示:

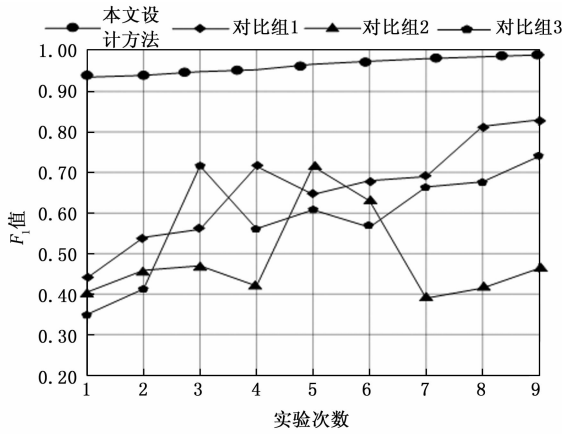
$$F_1 = \frac{2 * (U * R)}{(U + R)}$$

(7)

其中: U 为识别模型输出精确度, R 为识别模型输出召回率。

根据计算得到的结果展开实验测试。实验中,利用 4 种方法对多个数据集进行实体计算,统计其计算结果 F_1 值,得到具体统计结果如图 6 所示。

如图 6 所示,本文设计的方法在命名实体识别上取得了极高的 F_1 值,平均达到 0.92。这一卓越表现主要得益于大语言模型的优势,它能够捕捉文本中的深层语义信息,提高识别的准确性;同时,针对社交媒体中文文本的特性进行了优化,使得模型更能适应复杂多变的文本环境;相比之下,对比组的 3 种方法在 F_1 值方面均表现较差,平均 F_1 值分别为 0.65、0.52 和 0.48。

图 6 4 种方法的 F_1 值对比分析

这可能是因为这些方法在模型设计、技术选择或数据集使用等方面存在不足, 导致识别精度和召回率均较低。通过对比分析, 可以得出本文方法在命名实体识别任务中的显著优势。

3 结束语

本文通过预训练模型对海量社交媒体数据进行学习, 进而构建出一个对中文命名实体具有高度敏感性的识别方法。当面对新的社交媒体文本时, 该系统能够迅速而准确地识别出相应的中文实体, 从而为后续的信息抽取、情感分析、舆情监测等任务提供有力支持。该方法充分利用了大语言模型在处理自然语言文本时的强大能力, 特别是其深度学习和自我优化机制, 使其在处理复杂多变的社交媒体文本时表现出色。尽管本文方法在中文命名实体识别方面取得了显著成果, 但在领域适应性、实时性与效率方面仍存在一些挑战和局限性。在未来的研究将致力于克服这些挑战, 推动命名实体识别技术的不断发展和创新。

参考文献:

- [1] CHEN M, LUO X, YUAN P Y. A Chinese nested named entity recognition approach using sequence labeling [J]. International Journal of Web Information Systems, 2023, 19 (1): 42–60.
- [2] ZHANG Y, ZHANG Y, WANG S J, et al. A more cost-efficient Chinese named entity recognition based on trigger and matching network [J]. Journal of Intelligent & Fuzzy Systems; Applications in Engineering and Technology, 2023, 44 (2): 2085–2096.
- [3] KE X, WU X, OU Z, et al. Chinese named entity Recognition method based on multi-feature fusion and biaffine [J]. Complex & Intelligent Systems, 2024, 10 (5): 6305–6318.
- [4] GUO Q, GUO Y. Lexicon enhanced Chinese named entity recognition with pointer network [J]. Neural Computing and Applications, 2022, 34 (17): 14535–14555.

- [5] ZHANG L, XIA P, MA X, et al. Enhanced Chinese named entity recognition with multi-granularity BERT adapter and efficient global pointer [J]. Complex & Intelligent Systems, 2024, 10 (3): 4473–4491.
- [6] ZHANG T, SUN Y, WEN D U, et al. Judicial named entity recognition enhanced with semantic and boundary [J]. Journal of Tsinghua University (Science and Technology), 2024, 64 (5): 749–759.
- [7] XIANG Y, LIU W, GUO J Z L. Local and global character representation enhanced model for Chinese medical named entity recognition [J]. Journal of Intelligent & Fuzzy Systems; Applications in Engineering and Technology, 2023, 45 (3): 3779–3790.
- [8] CUI X, SONG C, LI D, et al. RoBGP: A Chinese nested biomedical named entity recognition model based on RoBERTa and global pointer [J]. Computers, Materials & Continua, 2024, 78 (3): 3603–3618.
- [9] WANG Y, LU L, YANG W, et al. Local or global? a novel transformer for Chinese named entity recognition based on multi-view and sliding attention [J]. International Journal of Machine Learning and Cybernetics, 2024, 15 (6): 2199–2208.
- [10] PAN J, ZHANG C, WANG H, et al. A comparative study of Chinese named entity recognition with different segment representations [J]. Applied Intelligence, 2022, 52 (11): 12457–12469.
- [11] TANG X, HUANG Y, XIA M, et al. A Multi-Task BERT-BiLSTM-AM-CRF strategy for Chinese named entity recognition [J]. Neural Processing Letters, 2022, 55 (2): 1209–1229.
- [12] NIE Y, ZHANG Y, PENG Y, et al. Borrowing wisdom from world: modeling rich external knowledge for Chinese named entity recognition [J]. Neural Computing and Applications, 2022, 34 (6): 4905–4922.
- [13] DIAO J, ZHOU Z, SHI G. Leveraging integrated learning for open-domain Chinese named entity recognition [J]. International Journal of Crowd Science, 2022, 6 (2): 74–79.
- [14] CHEN P, ZHANG M, YU X, et al. Named entity recognition of Chinese electronic medical records based on a hybrid neural network and medical MC-BERT [J]. BMC Medical Informatics and Decision Making, 2022, 22 (1): 1–13.
- [15] ZHANG Q, XUE C, SU X, et al. Named entity recognition for Chinese construction documents based on conditional random field [J]. Frontiers of Engineering Management, 2022, 10 (2): 237–249.
- [16] ZHANG Q, WU M, LV M Y H. Research on named entity recognition of Chinese electronic medical records based

- on multi-head attention mechanism and character-word information fusion [J]. *Journal of Intelligent & Fuzzy Systems: Applications in Engineering and Technology*, 2022, 42 (4): 4105–4116.
- [17] CHEN H, DAN L, LU Y, et al. An improved data augmentation approach and its application in medical named entity recognition [J]. *BMC Medical Informatics and Decision Making*, 2024, 24 (1): 1–13.
- [18] LIU H, PENG W, SONG J. RepEKShot: an evidential k-nearest neighbor classifier with repulsion loss for few-shot named entity recognition [J]. *The Journal of Supercomputing*, 2024, 80 (15): 22069–22098.
- [19] LEI X, SONG W, FAN R, et al. Semi-supervised geological disasters named entity recognition using few labeled data [J]. *GeoInformatica*, 2022, 27 (2): 263–288.
- [20] MA L L, YANG J, AN B, et al. Medical named entity recognition using weakly supervised learning [J]. *Cognitive*
- (上接第222页)
- [4] 孙 聪, 贾萌娜, 于起峰. 基于仿射近似投影模型的无人机对地目标定位方法 [J]. *中国惯性技术学报*, 2022, 30 (1): 104–112.
- [5] 张 赫, 乔 川, 匡海鹏. 基于激光测距的机载光电成像系统目标定位 [J]. *光学精密工程*, 2019, 27 (1): 8–16.
- [6] 郑 锴, 郑献民, 殷少锋, 等. 基于总体最小二乘的无人机实时目标定位方法 [J]. *电光与控制*, 2019, 26 (10): 26–29.
- [7] 张 岩, 李建增, 李德良, 等. 基于景象匹配的无人机侦察视频快速配准方法 [J]. *电光与控制*, 2017, 24 (5): 30–35.
- [8] 赵江洪, 刘茈菱, 杨 甲, 等. 几何单目视觉测距研究综述 [J]. *测绘科学*, 2023, 48 (9): 49–65.
- [9] 魏明鑫, 黄 浩, 胡永明, 等. 基于深度学习的多旋翼无人机单目视觉目标定位追踪方法 [J]. *计算机测量与控制*, 2020, 28 (4): 156–160.
- [10] LIU Y F, NOGUCHI N, LIANG L. Development of a positioning system using UAV-based computer vision for an airboat navigation in paddy field [J]. *Computers and Electronics in Agriculture*, 2019, 162: 126–133.
- [11] KIM D, LIU M, LEE S, et al. Remote proximity monitoring between mobile construction resources using camera-mounted UAVs [J]. *Automation in Construction*, 2019, 99: 168–182.
- [12] 徐 超, 高 敏, 曹 欢. 单目图像中姿态角估计的坦克目标测距方法 [J]. *光子学报*, 2015, 44 (5): 140–147.
- [13] LI Q, NIAZU, MIERIALDO B. An improved algorithm
- on Viola-Jones object detector [C] // Annecy, France; 2012 10th International Workshop on Content-Based Multimedia Indexing (CBMI), 2012: 1–6.
- [14] DALALN, TRIGGS B. Histograms of oriented gradients for human detection [C] // San Diego, CA, USA; 2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR05), 2005: 886–893.
- [15] FELZENSZWALB P, MCALLESTERD, RAMANAN D. A discriminatively trained, multiscale, deformable part model [C] // Anchorage, AK, USA; 2008 IEEE Conference on Computer Vision and Pattern Recognition, 2008: 1–8.
- [16] 江 波, 屈若锟, 李彦冬, 等. 基于深度学习的无人机航拍目标检测研究综述 [J]. *航空学报*, 2021, 42 (4): 137–151.
- [17] REDMON J, DIVVALA S, GIRSHICKR, et al. You only look once: unified, real-time object detection [C] // Las Vegas, NV, USA; 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2016: 779–788.
- [18] 周晓彦, 王 珂, 李凌燕. 基于深度学习的目标检测算法综述 [J]. *电子测量技术*, 2017, 40 (11): 89–93.
- [19] 李宏伟, 陈 冰, 张贺磊. 计算机视觉测距技术研究 [J]. *电视技术*, 2020, 44 (8): 60–61.
- [20] GRIFFIN B A, CORSO J J. Depth from camera motion and object detection [C] // Nashville, TN, USA; 2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 2021: 1397–1406.
- [21] 张华海, 郑南山, 王 军, 等. 由空间直角坐标计算大地坐标的简便公式 [J]. *全球定位系统*, 2002 (4): 9–12.