



基于弱依赖信息的知识库问答方法

吴天波¹, 刘露平¹, 罗晓东¹, 卿粼波^{1,2}, 何小海¹

(1. 四川大学 电子信息学院, 成都 610065; 2. 无线能量传输教育部重点实验室, 成都 610065)

摘要: 传统自动问答方法通常依赖谓词等先验信息实现知识库问答, 需要耗费较多的人力且泛化能力不佳。提出一种针对弱依赖信息的知识库问答方法, 结合 BERT 与 BiLSTM-CRF 网络提取问句中的命名实体, 定位知识库中与该实体相关的三元组信息, 通过答案匹配网络为三元组集合中的答案标上相似度分数, 使用阈值选择策略选取符合要求的答案集合, 并按照相似度分数由高到低排序后呈现给用户。实验结果表明, 该方法弱化了先验信息的依赖, 在减少人工干预的同时保证了问答质量, 并且在 NLPCC-ICCPOL-2016KBQA 数据集上取得了 87.05% 的 F1 分数。

关键词: 弱依赖信息; 知识库问答; 命名实体识别; 答案匹配; 阈值选择

开放科学(资源服务)标志码(OSID):



中文引用格式: 吴天波, 刘露平, 罗晓东, 等. 基于弱依赖信息的知识库问答方法[J]. 计算机工程, 2021, 47(6): 76-82.
英文引用格式: WU Tianbo, LIU Luping, LUO Xiaodong, et al. Knowledge base question answering method based on weak dependency information[J]. Computer Engineering, 2021, 47(6): 76-82.

Knowledge Base Question Answering Method Based on Weak Dependency Information

WU Tianbo¹, LIU Luping¹, LUO Xiaodong¹, QING Linbo^{1,2}, HE Xiaohai¹

(1. College of Electronics and Information Engineering, Sichuan University, Chengdu 610065, China;

2. Key Laboratory of Wireless Power Transmission, Ministry of Education, Chengdu 610065, China)

[Abstract] Traditional automatic question answering methods mostly rely on priori information such as predicate to realize knowledge base question answering, which is labor-intensive and leads to poor generalization performance. To address the problem, this paper proposes a KBQA method for weak dependency information. The method uses BERT in combination with the BiLSTM-CRF network to extract the named entity in a question. Then, it locates the triple information related to the entity in the knowledge base, and uses the answer matching network to give similarity scores to the answers in the triple set. Finally, it uses the threshold selection strategy to select the answer set that meets the requirements, and the answers are sorted according to the similarity before being presented to the user. Experimental results show that the method weakens the dependence on priori information, and reduces the manual intervention while ensuring the quality of KBQA. It achieves an F1 score of 87.05% on the NLPCC-ICCPOL-2016KBQA data set.

[Key words] weak dependency information; knowledge base question answering; named entity recognition; answer matching; threshold selection

DOI: 10.19678/j.issn.1000-3428.0058312

0 概述

自动问答是指利用计算机自动回答用户所提出的问题以满足用户知识需求的任务, 按照数据来源分为检索式问答、社区问答和知识库问答^[1]。知识库结构为 $\{E, A, V\}$ 三元组集合, 其中, E 表示实体, A

表示属性, V 表示目标值。知识库问答的核心目标是定位出问题所对应答案的三元组, 即所需的答案。例如, 在知识库中存在三元组 (“辣子鸡”, “分类”, “湘菜, 辣菜”), 当被问及 “辣子鸡属于什么菜系?” 时, 知识库问答将定位到该条三元组, 给出 “湘菜, 辣菜” 的答案。

基金项目: 国家自然科学基金(61871278); 四川省科技计划项目(2018HH0143); 成都市产业集群协同创新项目(2016-XT00-00015-GX)。

作者简介: 吴天波(1996—), 男, 硕士研究生, 主研方向为自然语言处理; 刘露平、罗晓东, 博士研究生; 卿粼波, 副教授; 何小海(通信作者), 教授。

收稿日期: 2020-05-13 **修回日期:** 2020-06-16 **E-mail:** nic5602@scu.edu.cn

知识库问答主要包括语义解析、信息抽取和向量建模3种途径。语义解析将自然语言转化为逻辑形式进行分析,使机器可以理解其中的语义信息,并从知识库中提取信息进行回答。文献[2]采用知识库问答方式,通过无监督手段将自然语言用解释器解析为逻辑形式,并在知识库中检索答案。信息抽取采用模糊检索方式,从问句中抽取关键信息,并以该信息为目标在知识库中检索更小的集合,在此集合上进一步得出答案。文献[3]对问题进行命名实体识别,利用实体信息从知识库中建立图模型,实现信息提取和答案筛选。向量建模将问题和答案映射到向量空间进行分析,近年来得益于深度学习的飞速发展,向量建模方法得到了广泛应用。文献[4]基于深度结构化语义模型匹配问题和谓语。文献[5]在图表示学习的基础上进行改进,从特定问题的子图中提取答案。文献[6]提出基于深度强化学习的网络,对问题和选项进行编码。文献[7]利用知识图嵌入将谓词和实体用低维向量表示,探索其在知识图谱问答任务中的潜在用途。

在中文领域,知识库问答多数结合信息抽取、向量建模两种方法实现。文献[8]利用深度卷积神经网络(Recurrent Neural Network, RNN)挖掘语义特征,并通过答案重排确定结果。文献[9]使用基于注意力机制的长短期记忆(Long Short-Term Memory, LSTM)网络^[10]将中文语料映射到向量空间,并依托实体抽取检索出备选的知识集合。文献[11]在此基础上引入人工规则,并结合句法分析进行关系词提取。文献[12]提出基于依赖结构的语义关系识别方法,从问句中挖掘深层的语义信息。上述方法都高度依赖问答对以外的信息,由于训练数据包括原始问答对以及每对问句和答案对应的三元组信息,因此这些信息在许多场景中并不具备,需通过大量人工标注或先验规则获得,耗费较多人力且泛化能力不佳,同时通常需要不同的预处理方法处理不同领域的问答数据。为解决上述问题,文献[13]提出非监督学习方法,利用动态规划思想,寻找全局最优决策,但问答结果的准确率不高。本文基于弱依赖信息,在仅已知问答对信息的情况下设计答案匹配策略,通过挖掘问句与答案潜在的语义联系以提高问答效率。

1 相关工作

1.1 整体流程

在知识库问答中,弱依赖信息是指数据来源仅含知识库和问答对,使得问答模型能尽可能少地依赖其他先验信息。基于弱依赖信息的知识库问答分为命名实体识别、答案匹配和阈值选择三大模块。首先通过命名实体识别提取问句中的实体,然后以

该实体为搜索条件生成查询语句,通过知识库检索返回三元组集合,并将去掉命名实体的问句与三元组集合中的答案集合依次做语义匹配,得到带相似度分数的一系列备选答案,最后通过阈值选择得出最终的答案。知识库问答整体流程如图1所示。

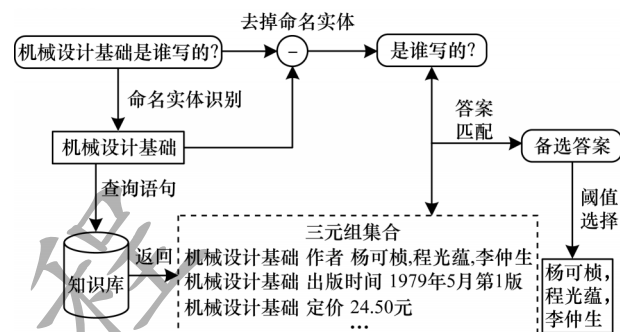


图1 知识库问答整体流程

Fig.1 Overall process of knowledge base question answering

在图1中,命名实体识别和答案匹配网络模型均使用BERT(Bidirectional Encoder Representations from Transformer)预训练模型进行特征提取。BERT模型内部使用Transformer代替卷积神经网络,能方便地迁移到其他网络中,输入的自然语言通过该基础网络后得到向量化的特征,再利用后续的网络结构实现各自功能。

1.2 BERT模型

BERT^[14]是Google AI团队于2018年提出的自然语言处理(Natural Language Processing, NLP)领域的通用模型,在信息抽取、语义推理和问答系统等众多任务中均取得了突破性的进展。BERT模型内部主要使用双向Transformer编码器,核心结构如图2所示。网络使用带注意力机制的双向Transformer block进行连接^[15],能更好地挖掘输入语料的上下文语义信息。

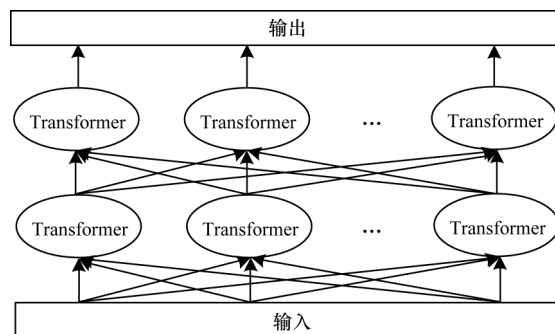


图2 BERT核心结构

Fig.2 Core structure of BERT

1.3 Transformer模型

Transformer^[16]是Google于2017年提出的基于

注意力机制的NLP经典模型。该模型分为编码器和解码器两部分,其中编码器结构如图3所示。Transformer通过在网络中引入多头注意力机制,调整输入的每个词的权重,因此能够获得更加全局的词向量表示。

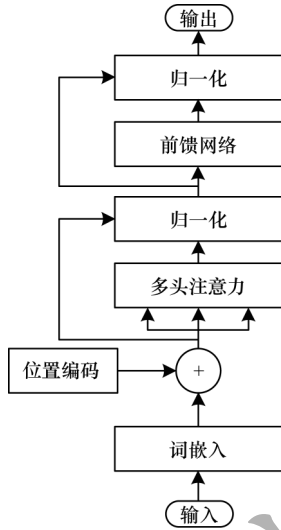


图3 Transformer编码器结构

Fig.3 Transformer encoder structure

2 知识库问答模型

2.1 命名实体识别

知识库问答的核心是使用命名实体识别算法提取出问句中的实体。命名实体识别是自然语言处理中的经典任务,它属于序列标注的子任务,通过对输入文字每个位置标注出相应的实体信息,实现实体抽取功能。实体标注有 BIO 和 BIOES 两种模式,本文采用 BIO 模式标注实体,其中,B-X 表示 X 实体的开头,I-X 表示 X 实体的中间或结尾,O 表示不是实体内容。由于本文研究的知识库问答的问句中仅涉及单一实体,因此仅定义一种实体类型 ENT。例如,当输入问句“李明的出生地是哪?”时,实体标注结果如图4所示。

李 明 的 出 生 地 是 哪 ?
B-ENT I-ENT O O O O O O O

图4 实体标注结果

Fig.4 Result of entity annotation

命名实体识别网络模型如图5所示,主要包括特征提取和实体标注两部分。在特征提取过程中,长度为 m 的输入问句被分割成词的序列 $\{w_1, w_2, \dots, w_m\}$ 送入BERT网络中,经分词及词嵌入后得到 m 个词向量。词向量经过 N 层的Transformer编码器特征提取后,得到长为序列长度 m 、宽为隐藏层维度 d 的特征矩阵,从而完成特征提取工作。在实体标注过程中通常采用BiLSTM-CRF^[17]网络进行重叠命名实体识

别^[18]。首先将特征矩阵输入到每个方向的神经元个数为 n 的双向LSTM层,进一步提取上下文的语义关联信息,其中, f 、 b 、 c 分别表示正向、逆向和输出神经元,输出的新特征向量隐藏层维度为 $2n$ 。该特征向量经过一层前馈神经网络,通过线性变换得到长度为 m 、宽度为待标注类型数的向量并将其作为CRF层的输入。

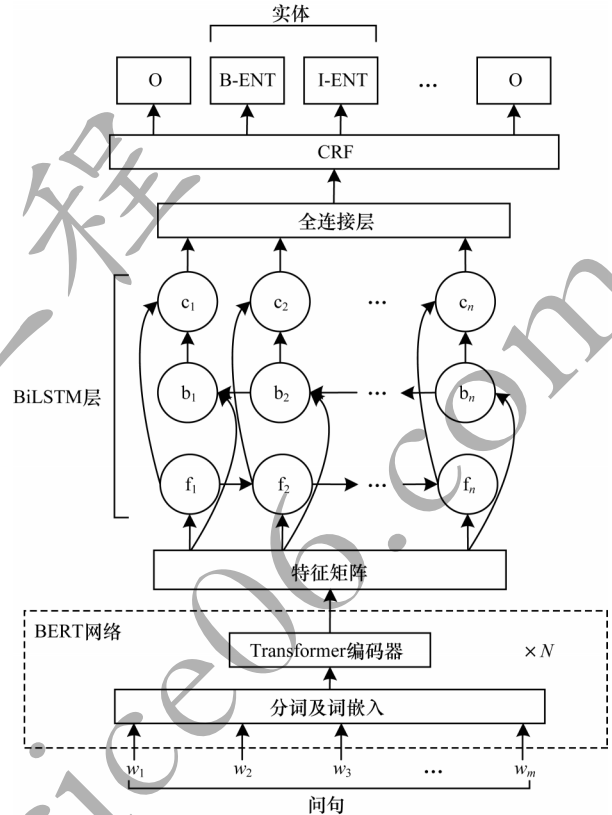


图5 命名实体识别网络模型

Fig.5 Network model of named entity recognition

由于本文仅定义一种实体类型,因此该向量宽度为3,分别代表B、I和O的状态分数。在CRF层中,线性链条件随机场概率模型对输入特征序列求出条件概率最大的输出标注序列,即为输入问句的每个位置标上标注信息。通过对输出标注序列的统计,便能定位出实体的起止位置。

在BiLSTM-CRF网络中,对于输入向量 x ,对应的输出为 y ,其得分计算如式(1)所示:

$$S(x, y) = \sum_i h_i[y_i] + P[y_{i-1}, y_i] \quad (1)$$

其中, h 表示BiLSTM层输出的三维向量, P 表示转移特征矩阵, $P[y_{i-1}, y_i]$ 表示输出标签从 y_{i-1} 到 y_i 的转移得分值。损失函数采用对数似然函数,训练时最小化式(2)中的目标函数:

$$L = \ln \sum_{y'} e^{S(x, y')} - S(x, y) \quad (2)$$

由于本文针对的数据集是单跳问答对,问句中抽取出的实体多为单个,如果存在多个实体,第一个

实体通常是问题的主语,因此将选取其作为候选实体。

2.2 答案匹配

在完成命名实体识别后,将提取的实体名作为关键词,生成知识库的查询语句,在知识库中检索返回包含该实体的三元组集合,为答案匹配做准备。在中文知识库问答中,通常将问句与三元组中的谓词做语义匹配,但这需要训练数据中包含的原始问答对以及具体的三元组信息,而特定任务的问答数据集通常没有这些额外信息,因此需要大量的人工标注或者特殊的预处理方式。本文提出的答案匹配方法直接将问句与答案信息做匹配,在训练时仅依赖原始问答对数据,在问答时计算知识库中三元组的答案与问句的匹配程度。首先对问句做预处理,去除命名实体,以防问句过长及冗余信息对答案匹配的效果产生干扰,然后将预处理后的问句与三元组集合中的每一个答案做相似度匹配,为每一个答案都标上相似度分数。相似度分数是一个0到1之间的值,因此在训练过程中,若输入为正确答案,则对应的相似度分数的标签为1,否则相似度分数的标签为0。

答案匹配网络模型如图6所示。问答对以[CLS]记号为开始,在每一次匹配中,预处理后的问句与答案之间用[SEP]记号隔开,连接成一个序列。

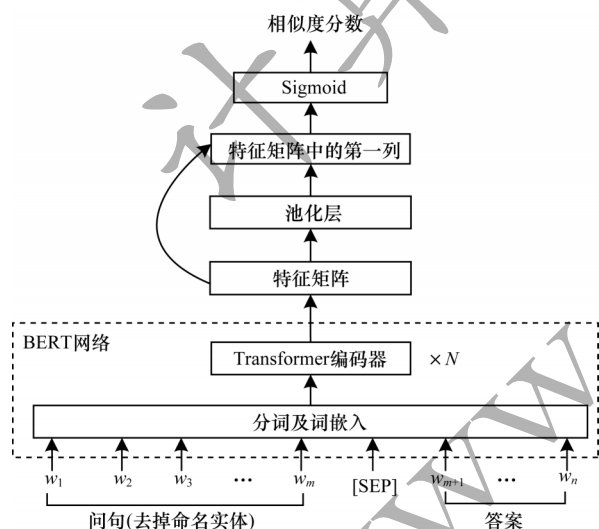


图6 答案匹配网络模型

Fig.6 Network model of answer matching

答案匹配网络的特征提取过程与命名实体识别网络类似,经过BERT网络后得到一个长为 $(m+n)$ 、宽为 d 的特征矩阵。由于网络最后一层为Sigmoid层,是分类网络的典型输出层,因此需要对特征矩阵进行下采样,使用一层池化层提取特征矩阵中最重要的信息,将特征矩阵的第一列(长为 d)提取出来,作为Sigmoid层的输入。最终经Sigmoid层输出,得到一个0到1之间的值,即相似度分数。

由于答案匹配网络的最后一层为Sigmoid层,因此损失函数采用交叉熵损失函数。标签仅有0和1两类,损失函数结构与二分类任务中的结构类似。在一次相似度匹配中,若样本标签为 y ,则预测的相似度分数为 s ,损失函数表示为:

$$L = -[y \cdot \ln s + (1 - y) \cdot \ln(1 - s)] \quad (3)$$

2.3 阈值选择

通过答案匹配为包含问句中实体的三元组集合的每一个答案都标上相似度分数,之后基于这些相似度分数选出合适的答案。较简单的做法是选出相似度分数最高的答案,这在基于谓词匹配的传统方法中具有最好的效果,但将其使用在本文提出的答案匹配的方法中得出的答案会有一定误差,这是由于答案匹配得到的相似度分数通常比谓词匹配小很多,因此相近的答案之间区分度不高。

知识库问答的评测指标主要为F1分数(F),假定标准答案和预测答案均为集合形式,通过精确率(P)和召回率(R)计算得到F1分数。精确率表示预测正确的答案在预测答案集合中所占的比例,反映了问答系统的准确程度。召回率表示预测正确的答案在正确答案集合中所占的比例,反映了问答系统的完备程度。一个高质量的问答系统应该同时保持高的精确率值和召回率值,并通过F1分数对其性能进行评价。F1分数的计算公式为:

$$F = \frac{2 \times P \times R}{P + R} \quad (4)$$

若要构建性能良好的问答系统,只有在答案选择中返回相似度分数近似的答案集合,并将预测答案的错误和遗漏同时控制到最低,才能得到较高的F1分数。本文采用阈值选择策略,通过实验对比选择合适的相似度阈值,高于阈值的答案将被选中,构成预测答案的集合,并按相似度分数的高低排序后呈现给用户。使用 S 表示每个问题的相似度分数, $S_{\text{threshold}}$ 表示设定的相似度阈值,每个答案的选中状态为 B , $B=1$ 表示答案被选中, $B=0$ 表示答案未被选中,计算公式为:

$$B = \begin{cases} 1, & S \geq S_{\text{threshold}} \\ 0, & S < S_{\text{threshold}} \end{cases} \quad (5)$$

3 实验与结果分析

3.1 实验数据集与环境

本文使用NLPCC-ICCPOL-2016KBQA数据集发布的知识库和问答对数据,共有14 609个训练问答对和9 870个测试问答对。为使评测结果更加客观,进一步将训练问答对随机划分为训练集和开发集,测试问答对作为测试集。数据集划分情况如表1所示。

表1 数据集划分情况

Table 1 Division of dataset

名称	数量
训练集	12 000
开发集	2 609
测试集	9 870

本文实验运行在CPU为Inter i5-4590、内存为12 GB的计算机上,模型训练所用显卡为Nvidia GTX 1080Ti,显存为11 GB,所用深度学习框架为CUDA 10.0和Tensorflow 1.14,操作系统为64位Windows 10,知识库数据存储和检索使用Mysql 5.6.46。

3.2 命名实体识别结果分析

知识库问答中的问句格式较为固定,任何一个短句单实体的数据集都可以作为命名实体识别的训练数据。本文为了验证实验结果,所用数据为问答对中的问题及其所含的实体信息,使用的BERT模型为中文版本,通过对加载的预训练参数进行微调的方式,在12 000个训练集问题上对命名实体识别网络模型进行训练,并分别在训练集、开发集和测试集上进行性能测试。命名实体识别的超参数设置如表2所示。

表2 命名实体识别的超参数设置

Table 2 Hyperparameters setting of named entity recognition

名称	取值
最大序列长度	64
Transformer头数量	12
Transformer层数	12
隐藏层维度	768
LSTM神经元个数	128
batch尺寸	32
学习率	5e-5
总迭代次数	20
dropout比率	0.5

训练过程总共迭代7 028次,采用带权值衰减的Adam优化器^[19]优化损失函数。在训练集上完成训练后,分别将模型在训练集、开发集、测试集上进行性能测试,结果如表3所示。限于语料库规模,测试结果存在轻微的过拟合现象,在训练集上基本准确,在开发集和测试集上有一定误差,总体表现较好。

表3 命名实体识别测试结果

Table 3 Test results of named entity recognition %

数据集	准确率	精确率	召回率	F1分数
训练集	100.00	99.99	100.00	100.00
开发集	99.57	96.95	97.23	97.09
测试集	99.53	97.05	97.26	97.15

3.3 答案匹配结果分析

为训练答案匹配的网络模型,需要在已有问答对的基础上制作答案匹配的数据集。具体地,将每个问题去掉命名实体后,与答案相连接,再在后面加上一个“1”,表示该问句与答案的相似度为1,连接处均用[SEP]标记隔开。制作负样本的过程与问答流程类似,以命名实体为关键词在知识库中检索,得到与该实体有关的答案集合,将不为该问题答案的名词以同样的方式连接在问句后,并在后面加上一个“0”,表示问句与该答案的相似度为0。对于知识库中仅有一个三元组的实体,为加以区分,则在以其他实体为关键字的三元组中随机选取5个答案作为负样本,添加到数据集中,得到的答案匹配数据集规模如表4所示。

表4 答案匹配数据集规模

Table 4 Dataset size of answer matching

数据集	正样本数量	负样本数量	总数量
训练集	12 000	63 109	75 109
开发集	2 609	37 696	40 305
测试集	9 870	46 346	56 216

将训练集数据输入答案匹配网络进行训练。由于网络特征提取部分同样使用BERT,因此超参数选取除了没有LSTM以外,其他设置与命名实体识别一致。模型训练的优化器同样采用带权值衰减的Adam,网络共迭代11 505次。由于相似度分数为0到1之间的值,不可能与标签完全相等,在计算测试指标时,将网络输出修改为类别,即将其当作一个二分类问题,只能输出“0”或“1”。在计算性能指标时,除了准确率以外,AUC也是一个重要的性能指标,它能够更客观地衡量模型对答案匹配数据集的分类效果。答案匹配模型在训练集、开发集和测试集上的测试结果如表5所示。由于数据规模有限,本文模型在开发集和测试集上的表现较差,但AUC值均达到86%以上,为最终的自动问答质量提供了保障。

表5 答案匹配测试结果

Table 5 Test results of answer matching %

数据集	准确率	AUC
训练集	99.41	98.64
开发集	96.46	87.55
测试集	92.58	86.27

3.4 阈值选择结果分析

在完成命名实体识别和答案匹配模型的训练后即可进行知识库问答。在未加入阈值选择机制时,直接选择知识库中包含实体的三元组集合中相似度分数最高的答案作为输出,得到的问答结果如表6

所示。由于标准答案和预测答案均仅有一个,准确率、精确率、召回率和F1分数为相同的值,因此仅列出了F1分数的结果。

表6 未加入阈值选择机制的问答结果

Table 6 Question answering results without threshold

selection mechanism		%
数据集	F1 分数	
训练集	95.28	
开发集	87.77	
测试集	86.79	

通过记录回答错误的问题,并对其相似度分数进行观察,发现除数据集本身存在的噪声外,较为模糊的答案中前几名之间的相似度分数都比较接近,主要在 10^{-5} 至 10^{-2} 附近。为确定在本文数据集下的最佳阈值,分别选取 10^{-2} 、 10^{-3} 、 10^{-4} 和 10^{-5} 这4个相似度阈值,在开发集上调用阈值选择机制进行测试,结果如表7所示。在具体执行时,对于具有最高相似度的答案仍低于阈值的情况,直接将该答案作为问题的输出。

表7 不同相似度阈值下的开发集问答结果

Table 7 Question answering results of development set with different similarity thresholds

相似度阈值	精确率	召回率	F1 分数
10^{-2}	86.37	89.61	87.96
10^{-3}	85.95	90.07	87.96
10^{-4}	85.61	90.95	88.20
10^{-5}	84.16	91.15	87.52

从阈值选择的结果可以看出,随着选取阈值的减小,精确率逐渐变小,召回率逐渐变大,这是备选答案增多带来的必然结果。当阈值为 10^{-4} 时,开发集上问答的F1分数最高;当阈值进一步下降时,精确率因选中的答案过多而下降较多,因此F1分数也随之降低。

3.5 自动问答结果对比

通过阈值选择的实验结果,选取 10^{-4} 为本文知识库问答的相似度阈值,将其应用在最终的问答系统中,测试结果如表8所示。训练集和开发集均来源于NLPCC-ICCPOL-2016KBQA任务原始问答对的训练集,在公开的评测指标中以测试集的F1分数为准。

表8 知识库问答最终结果

Table 8 Final results of knowledge base

question answering		%	
数据集	精确率	召回率	F1 分数
训练集	93.03	95.69	94.34
开发集	85.61	90.95	88.20
测试集	84.59	89.65	87.05

本文问答系统在实际应用中将阈值选择作为可选开关。在许多应用场景中,问答任务要求返回单一答案,此时将关闭阈值选择开关,将相似度最高的答案呈现给用户。若用户对答案有疑惑,或者一些场景允许返回多个答案,则可以开启阈值选择,将候选答案集按相似度从高到低的顺序呈现。

本文选取DPQA^[13]、NEU(NLP Lab)、HIT-SCIR、CCNU、InsunKBQA^[9]、NUDT、PKU^[8]、WHUT^[11]和WenRichard^[20]作为对比方法,自动问答结果如表9所示。DPQA基于动态规划思想进行研究,其无监督思路具有参考意义,但问答效果较为受限。PKU、NUDT、CCNU、HIT-SCIR和NEU(NLP Lab)分别是NLPCC-ICCPOL-2016KBQA任务评测成绩的前5名的自动问答方法,它们主要依靠一些人工规则保证问答性能,例如PKU构造正则表达式以去除问句中的冗余信息,NUDT使用词性的组合特征实现命名实体识别等。InsunKBQA是基于知识库三元组中谓词的属性映射构建的自动问答方法,加入了少量人工特征。WHUT是通过句法分析等方式实现的自动问答方法。WenRichard首先在NLPCC-ICCPOL-2016KBQA数据集上应用BERT进行特征提取,并取得了目前公开的最好结果。本文方法除了应用BERT,还对答案选择方法进行改进,将其分解为答案匹配和阈值选择两个步骤,减少了对人工标注和预处理的需求,得到的测试集F1分数为87.05%,具有最优的性能表现。

表9 10种方法的自动问答结果

Table 9 Automatic question answering results of ten methods

自动问答方法	测试集 F1 分数
DPQA	71.00
NEU(NLP Lab)	72.72
HIT-SCIR	79.14
CCNU	79.57
InsunKBQA	81.35
NUDT	81.59
PKU	82.47
WHUT	82.94
WenRichard	84.71
本文方法	87.05

4 结束语

本文针对弱依赖信息,提出一种基于问答对数据的知识库自动问答方法。通过命名实体识别网络提取问句中的实体,同时以该实体名为关键词获取相关三元组集合,利用答案匹配网络为每一个答案标注相似度分数,最终通过阈值选择筛选备选答案并输出结果。实验结果表明,知识库问答方法在NLPCC-ICCPOL-2016KBQA数据集上的F1分数为

87.05%,其中的答案选择方法弱化了对问答数据中谓词等先验信息的依赖,无需人工干预就能在一个问答对数据集上完成训练,具有良好的泛化性能。通过实验发现本文知识库问答方法对数字类型的答案筛选精度有待提高,后续将利用表示学习等方法从候选答案集合中筛选出最优答案,进一步提升问答质量。

参考文献

- [1] DIEFENBACH D, LOPEZ V, SINGH K, et al. Core techniques of question answering systems over knowledge bases: a survey[J]. Knowledge and Information Systems, 2018, 55(3): 529-569.
- [2] BERANT J, CHOU A, FROSTIG R, et al. Semantic parsing on freebase from question-answer pairs [C]// Proceedings of 2013 Conference on Empirical Methods in Natural Language Processing. Philadelphia, USA: ACL Press, 2013: 1533-1544.
- [3] YAO X, DURME B V. Information extraction over structured data: question answering with freebase [C]// Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics. Philadelphia, USA: ACL Press, 2014: 956-966.
- [4] XIE Zhiwen, ZENG Zhao, ZHOU Guangyou, et al. Topic enhanced deep structured semantic models for knowledge base question answering [J]. Science China Information Sciences, 2017, 60(11): 28-42.
- [5] SUN H, DHINGRA B, ZAHEER M, et al. Open domain question answering using early fusion of knowledge bases and text [EB/OL]. [2020-04-02]. <https://arxiv.org/abs/1809.00782>.
- [6] YU Jianxing, ZHA Zhengjun, YIN Jian. Inferential machine comprehension: answering questions by recursively deducing the evidence chain from text [C]// Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics. Philadelphia, USA: ACL Press, 2019: 2241-2251.
- [7] HUANG Xiao, ZHANG Jingyuan, LI Dingcheng, et al. Knowledge graph embedding based question answering [C]// Proceedings of the 20th ACM International Conference on Web Search and Data Mining. New York, USA: ACM Press, 2019: 105-113.
- [8] LAI Yuxuan, JIA Yanyan, LIN Yang, et al. A Chinese question answering system for single-relation factoid questions [EB/OL]. [2020-04-02]. <https://arxiv.org/abs/1809.00782>.
- [9] ZHOU Botong, SUN Chengjie, LIN Lei, et al. LSTM based question answering for large scale knowledge base [J]. Journal of Peking University (Natural Science Edition), 2018, 54(2): 286-292. (in Chinese)
周博通, 孙承杰, 林磊, 等. 基于LSTM的大规模知识库自动问答[J]. 北京大学学报(自然科学版), 2018, 54(2): 286-292.
- [10] HOCHREITER S, SCHMIDHUBER J. Long short-term memory [J]. Neural Computation, 1997, 9(8): 1735-1780.
- [11] ZHANG Fangrong, YANG Qing. Research on entity relation extraction method in knowledge-based question answering [J]. Computer Engineering and Applications, 2020, 56(11): 219-224. (in Chinese)
张芳容, 杨青. 知识库问答系统中实体关系抽取方法研究[J]. 计算机工程与应用, 2020, 56(11): 219-224.
- [12] DUAN Jiangli, HU Xin. Semantic relation recognition for natural language question answering [J]. Journal of Shandong University (Engineering Science), 2020, 50(3): 1-7. (in Chinese)
段江丽, 胡新. 自然语言问答中的语义关系识别[J]. 山东大学学报(工学版), 2020, 50(3): 1-7.
- [13] WANG Yue, ZHANG Richong. Question answering over knowledge base using dynamic programming [J]. Journal of Zhengzhou University (Science Edition), 2019, 51(4): 37-42. (in Chinese)
王玥, 张日崇. 基于动态规划的知识库问答方法[J]. 郑州大学学报(理学版), 2019, 51(4): 37-42.
- [14] DEVLIN J, CHANG M W, LEE K, et al. BERT: pre-training of deep bidirectional transformers for language understanding [C]// Proceedings of 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies. Philadelphia, USA: ACL Press, 2019: 4171-4186.
- [15] SCHUSTER M, PALIWAL K K. Bidirectional recurrent neural networks [J]. IEEE Transactions on Signal Processing, 1997, 45(11): 2673-2681.
- [16] VASWANI A, SHAZEER N, PARMAR N, et al. Attention is all you need [C]// Proceedings of the 31st International Conference on Neural Information Processing Systems. Philadelphia, USA: ACL Press, 2017: 5998-6008.
- [17] MA X, HOVY E. End-to-end sequence labeling via bidirectional LSTM-CNNs-CRF [C]// Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics. New York, USA: ACM Press, 2016: 1064-1074.
- [18] WEN Xiuxiu, MA Chao, GAO Yuanyuan, et al. A Chinese overlapping name entity recognition method based on label clustering [J]. Computer Engineering, 2020, 46(5): 41-46. (in Chinese)
温秀秀, 马超, 高原原, 等. 基于标签聚类的中文重叠命名实体识别方法[J]. 计算机工程, 2020, 46(5): 41-46.
- [19] LOSHCILLOV I, HUTTER F. Decoupled weight decay regularization [C]// Proceedings of International Conference on Learning Representations. Washington D. C., USA: IEEE Press, 2019: 1-8.
- [20] WenRichard. KBQA-BERT [EB/OL]. [2020-04-02]. <https://github.com/WenRichard/KBQA-BERT>.

编辑 陆燕菲