

基于关键词结构编码的涉案微博评价对象抽取模型

王静赞^{1,2}, 余正涛^{1,2}, 相艳^{1,2}, 陈龙^{1,2}

(1. 昆明理工大学信息工程与自动化学院, 昆明 650500; 2. 昆明理工大学云南省人工智能重点实验室, 昆明 650500)

摘要: 涉案微博评价对象抽取旨在从微博评论中识别出用户评价的案件对象词项, 有助于掌握大众对于特定案件不同方面的舆论。现有方法通常将评价对象抽取视为一个序列标注任务, 但并未考虑涉案微博的领域特点, 即评论通常围绕正文中出现的案件关键词展开讨论。为此, 本文提出一种基于关键词结构编码的序列标注模型, 进行涉案微博评价对象抽取。首先从微博正文中获取多个案件关键词, 并使用结构编码机制将其转换为关键词结构表征, 然后将该表征通过交互注意力机制融入评论句子表征, 最后利用条件随机场 (Conditional random field, CRF) 抽取评价对象词项。在两个案件的数据集上进行了实验, 结果表明: 相较于多个基线模型, 本文方法性能得以提升, 验证了所提方法的有效性。

关键词: 结构编码; 涉案微博; 舆情; 评价对象抽取

中图分类号: TP391 **文献标志码:** A

A Model for Extracting Evaluation Objects of Cased - Involved Microblog Based on Keyword Structured Encoding

WANG Jingyun^{1,2}, YU Zhengtao^{1,2}, XIANG Yan^{1,2}, CHEN Long^{1,2}

(1. Faculty of Information Engineering and Automation, Kunming University of Science and Technology, Kunming 650500, China;
2. Yunnan Key Laboratory of Artificial Intelligence, Kunming University of Science and Technology, Kunming 650500, China)

Abstract: The purpose of extracting evaluation object of the microblog involved in a case is to identify the case object terms of the user evaluation from the microblog comments, which helps to grasp public thought on different aspects of a certain case. In general, the existing methods regard evaluation object extraction as a sequence labeling task, but do not take into account the domain characteristics of the microblog involved in the case, that is, comments are usually discussed around the case keywords that appear in the microblog text. For this reason, this paper proposes a sequence labeling model based on case keyword structured encoding to extract the evaluation objects of the microblog involved in the case. First of all, a number of case keywords are obtained from the text of microblogs, and the structured encoding mechanism is used to convert them into keyword structural representations. After that, the representations are integrated into the comment sentence representation through the cross attention mechanism. In the end, the evaluation target terms are extracted by the conditional random field (CRF). Experiments are conducted on the data sets of

基金项目: 国家重点研发计划 (2018YFC0830105, 2018YFC0830101, 2018YFC0830100); 云南省重大科技专项计划项目 (202002AD080001); 云南省基础研究专项面上项目 (202001AT070047, 202001AT070046)。

收稿日期: 2021-08-30; **修订日期:** 2022-01-27

two cases. Compared with the multiple baselines, the encouraging progress validates the effectiveness of the proposed approach.

Key words: structured encoding; microblog involved in the case; public opinion; evaluation object extraction

引言

微博等社交媒体的蓬勃发展让人们获得更丰富、更及时的信息,同时每天也会产生大量评论。其中,与案件相关的评论在网络上迅速传播,所产生的舆论会干扰有关机构的工作。为此,获取涉案微博评论的评价对象,对于后续进行案件相关评论的细粒度情感分析,掌握案件舆论走向具有重要的作用。

涉案微博评价对象抽取的目的是从微博用户的评论文本中识别出被评价的对象,例如,在“这次事故女司机是无辜的”这一评论中,需要识别出评价对象“女司机”。针对评价对象抽取任务,目前的方法主要分为基于规则^[1-4]、基于传统机器学习^[5-8]和基于深度学习的方法^[9-17]。基于规则的抽取方法是利用关联规则、频繁词项和句法依存关系进行评价对象的抽取,但是这些方法严重依赖于人工定义的规则,需要消耗大量的时间精力,可移植性差。基于传统机器学习的方法通常将该任务看作一个序列标注任务,利用支持向量机、句法关系和条件随机场(Conditional random field, CRF)等进行评价对象的抽取,也依赖于手工特征。基于深度学习的模型可以自动地学习特征,通过注意力机制等获取句子中与评价对象更相关的信息,在评价对象抽取任务上表现出优异的性能。

涉案微博评价对象抽取任务具有特定的领域性。针对某个案件,网友的评论通常都会围绕微博正文中所提及的案件发生的人物、地点等关键词展开。换句话说,正文中出现的案件关键词构成了用户评论的评价对象。如表1所示,在“奔驰女司机维权案”这个案件中,评论中出现的“女司机”“奔驰店”“女车主”“奔驰”等评价对象在对应的正文中都被提及。显然,微博正文中与案件相关的关键词信息对于涉案微博评价对象抽取任务是有效的。

表1 微博正文及其对应的评论示例
Table 1 Example of microblog text and corresponding comments

微博正文	评论
4月9日,陕西西安一女车主,因坐奔驰车顶哭诉维权难走红,有消息称双方已和解。4月12日,女车主家人小磊(化名)称,双方并未和解。13日,涉事奔驰店家称,已经在和客户协商解决此事	(1)心疼女司机,给自己买个礼物还这么糟心; (2)这家奔驰店也太黑心了,支持女车主维权; (3)以后再也不买奔驰,真的是烂透了; (4)奔驰4s店是应该整顿一下了,猫腻太多了。

因此,可以利用微博正文中的案件关键词信息作为引导,帮助模型识别用户评论中的评价对象。然而,案件关键词之间是独立的,不存在序列依赖关系。因此,借助Lin等^[18]提出的结构编码思想,本文将关键词组合表示为多个语义片段,从而综合利用多个关键词的信息指导评价对象的抽取。具体来说,首先使用结构编码机制对案件关键词进行编码表征,然后通过交互注意力机制融合评论句子表征和案件关键词表征,最后送入CRF,进行评价对象的抽取。

本文的主要贡献如下:

- (1) 结合涉案微博数据的特点,提出了利用微博正文中的关键词信息指导评论中的评价对象抽取。
- (2) 提出利用结构编码机制对微博正文关键词进行编码,从而能综合利用多个关键词信息,构建涉案微博评价对象抽取模型。

(3) 在人工标注的“奔驰女司机维权案”和“重庆公交车坠江案”涉案微博数据集上进行实验,证明了模型的有效性。

1 相关工作

现有的评价对象抽取的模型主要分为基于规则、基于传统机器学习和基于深度学习的抽取模型。

(1) 基于规则的抽取方法是利用人工定义的规则和句法依存关系抽取评价对象词项。例如,Hu等^[1]提出利用关联规则抽取评价对象,并抽取出句子中与评价对象相距最近的表达意见的形容词;Zhuang等^[2]采用 WordNet 和人工构建的词表寻找评价对象和观点词,然后使用依存解析的方法提取有效的评价对象-观点对;Blair-Goldensohn等^[3]提出利用句子中出现的频繁词项抽取评价对象;Qiu等^[4]提出了一种使用依存树学习句法关系的独特方法。然而,这些方法严重依赖于预定义的规则和外部资源。

(2) 基于传统机器学习的抽取方法一般是将其定义为一个序列标注任务。例如,Kessler等^[5]提出使用支持向量机对情感表达的潜在评价对象进行排序,通过解析句子的依存关系抽取评价对象词项;Jin等^[6]提出构建一个由词和词性组成的词集,根据这个词集对评论进行标注,然后送入隐马尔可夫模型训练,抽取评价对象词项;Jakob等^[7]提出使用通用句法关系作为枢纽来减少领域间的差异,用于跨领域的评价对象抽取;Shu等^[8]提出了一种终身学习的方法能让传统的 CRF 从之前领域的抽取任务中获取知识,从而利用获取的知识在当前的任务中更好地执行抽取任务。然而,上述方法对大量的手工特征有很强的依赖性,难以处理特征缺失的情况。

(3) 基于深度学习的抽取方法可以自动学习特征,在该任务上表现出了强大性能。例如,Li等^[10]利用神经记忆交互建立评价对象和观点词之间的联系,并同时考虑了局部和全局的记忆进行评价对象抽取。为了更加关注句子中重要的词,Li等^[11]提出将之前时刻的评价对象预测的信息结合到当前时刻的评价对象中,生成有历史信息的评价对象。Wang等^[12]提出耦合多层注意力模型,挖掘评价对象之间间接的关联,从而得到更精准的评价对象信息。Ma等^[13]设计了门控单元网络,将对应的单词表示纳入解码器,使用位置感知注意力,更加关注目标单词的相邻单词。He等^[14]和 Luo等^[15]提出使用注意力机制的自动编码器进行评价对象抽取,并在各个数据集上取得了较好的性能。近期,Tulkens等^[16]提出了一种新的注意力机制,可以从句子中获取更多相关信息。为了解决缺少大量标注数据的问题,Li等^[17]提出一种新的数据增强方法,将数据增强看作一个条件增强任务,生成句子的同时对齐原始的标签序列,同时能够保留原始的评价对象不变,进行评价对象抽取。由于字符级信息对于中文文本的序列标注任务较为重要,所以 Li等^[9]提出使用字符级信息提升评价对象抽取模型的性能。

从上述研究工作来看,基于深度学习的模型对于评价对象抽取任务是有效的,但是对于涉案微博这类特定数据而言,有其自身的领域特性。因此,本文提出基于关键词结构编码的涉案微博评价对象抽取方法。

2 本文方法

本文提出的评价对象抽取模型使用 B(begin)、I(inside)和 O(other)进行字符级别的标签标注。其中,B代表评价对象的起始位置,I代表评价对象内部,O代表不是评价对象。数据标注示例如表 2 所示。

本文模型的基本结构分为评论编码层、关键词编码层、信息融合层和评价对象抽取层 4 部分,如图 1 所示。其中,评论编码层将微博评论句的字符嵌入和词嵌入送入 BiLSTM 进行编码,并将得到的编码表示进行拼接,输入双层高速网

表 2 数据标注示例

Table 2 Data annotation example

这	次	事	故	女	司	机	是	无	辜	的
O	O	O	O	B	I	I	O	O	O	O

络^[19];关键词编码层将案件关键词嵌入送入BiLSTM编码,然后通过结构编码机制进一步提取结构编码表征;信息融合层通过交互注意力机制将评论句子表征和案件关键词结构表征进行融合;评价对象抽取层将最终的特征表示送入CRF,抽取评价对象词项。

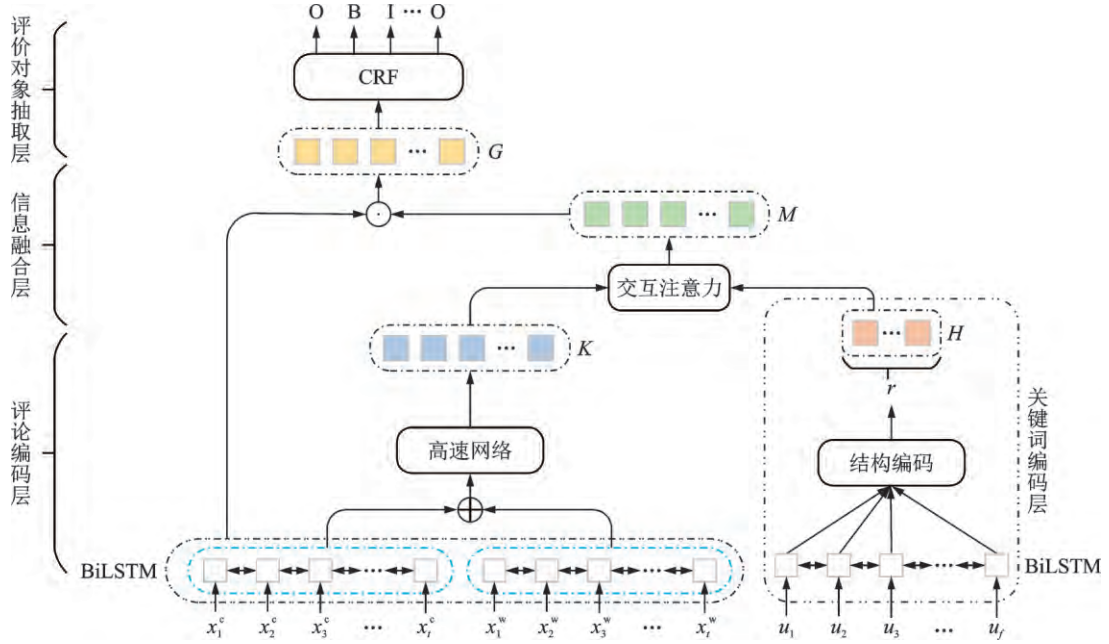


图1 涉案微博的评价对象抽取示意图

Fig.1 Schematic diagram of evaluation object extraction of the case-involved microblog

2.1 评论编码层

由于涉案微博评论的主观性和表达方式的自由性,其评价对象可能会出现多字少字的情况,造成分词结果不准确。因此本文按照字符粒度对评论数据集进行标注,并且同时使用字符嵌入和词嵌入作为评论的初始表示。给定一个微博评论句子,字符序列和词序列按照字符长度进行对齐。即对于每个字符,都为其分配包含该字符的词。字符向量序列和词向量序列分别表示为 $X^c = \{x_1^c, x_2^c, \dots, x_j^c, \dots, x_t^c\}$ 和 $X^w = \{x_1^w, x_2^w, \dots, x_j^w, \dots, x_t^w\}$, 其中 t 表示句子中字符的总个数, x_j^c 表示句子中第 j 个位置的字符, x_j^w 表示第 j 个位置的字符对应的词。将词嵌入和字符嵌入分别输入 BiLSTM 进行编码,并将编码得到的隐表示 $H^c \in \mathbb{R}^{2d \times t}$ 和 $H^w \in \mathbb{R}^{2d \times t}$ 进行拼接得到表示 $H^{cw} \in \mathbb{R}^{4d \times t}$, 具体如下

$$H^c = \text{BiLSTM}(X^c) \quad (1)$$

$$H^w = \text{BiLSTM}(X^w) \quad (2)$$

$$H^{cw} = H^c \oplus H^w \quad (3)$$

式中: $\oplus$ 表示拼接操作; d 表示嵌入维度。

然后将其输入双层高速网络^[19], 以此平衡字符向量和词向量的贡献比, 得到具有上下文语义特征的评论多层次向量表示 $K \in \mathbb{R}^{2d \times t}$ 。

$$K = O(H^{cw}, W_o) \cdot T(H^{cw}, W_T) + H^{cw} \cdot C(H^{cw}, W_C) \quad (4)$$

式中: O 表示非线性函数; T 表示转换门; C 表示携带门; W_o 、 W_T 和 W_C 为权重矩阵。

2.2 关键词编码层

本文使用 TextRank 分别从两个案件的微博正文中抽取案件关键词。给定一个评论句对应的一组案件关键词,其词向量序列表示为 $U = \{u_1, u_2, \dots, u_f\}$, 其中 f 表示关键词的总个数。将其送入 BiLSTM, 得到具有上下文语义特征的案件关键词向量表示 $L \in \mathbb{R}^{2d \times f}$ 。

$$L = \text{BiLSTM}(U) \quad (5)$$

本文使用结构编码操作将具有上下文语义特征的案件关键词的向量表示 $L \in \mathbb{R}^{2d \times f}$ 转化为结构化表示 $H \in \mathbb{R}^{2d \times r}$, 具体如下

$$A = \text{Softmax}(W_2 \tanh(W_1 L^T)) \quad (6)$$

$$H = AL \quad (7)$$

式中: $A \in \mathbb{R}^{r \times f}$ 为权重矩阵; W_1 和 W_2 为可训练的参数; r 为超参数, 表示 $L \in \mathbb{R}^{2d \times f}$ 转换为结构化表示的结构数量。

因为权重矩阵 $A \in \mathbb{R}^{r \times f}$ 容易面临过度平滑的问题, 即权重矩阵中的值趋于相似。为了让权重矩阵中的值更加具有区分性, 本文引入了惩罚项。具体而言, 受 Lin 等^[18] 的启发, 本文使用惩罚项 Z 作为损失函数中的一部分来保证 H 中结构化表示的多样性。

$$Z = \|AA^T - I\|_F^2 \quad (8)$$

式中: I 表示单位矩阵; $\|\cdot\|_F$ 为矩阵的 Frobenius 范数。

2.3 信息融合层

将具有上下文语义特征的评论多层次向量表示 $K = \{k_1, k_2, \dots, k_t\} \in \mathbb{R}^{2d \times t}$ 与关键词编码层得到的结构化表示 $H = \{h_1, h_2, \dots, h_r\} \in \mathbb{R}^{2d \times r}$ 做交互注意力, 由此得到的关键词表征 $M \in \mathbb{R}^{2d \times t}$ 用来表示评论句, 具体操作如下。

如图 2 所示的交互注意力机制, 对结构化表示 $H = \{h_1, h_2, \dots, h_r\} \in \mathbb{R}^{2d \times r}$ 中的每一个特征表示进行加权求和, 由此得到信息交互的关键词表征。

$$m_j = \sum_{i=1}^r \alpha_{j,i} \cdot h_i \quad (9)$$

式中注意力权重 $\alpha_{j,i}$ 用相应的匹配分数 $s_{j,i}$ 通过 softmax 函数计算得到, $s_{j,i}$ 通过特征向量 k_j 和 h_i 的双线性乘积计算得到

$$\alpha_{j,i} = \frac{e^{s_{j,i}}}{\sum_{x=1}^t e^{s_{j,x}}} \quad (10)$$

$$s_{j,i} = \tanh(k_j W h_i + b) \quad (11)$$

式中 W 和 b 为可训练的参数。

将上述融合之后的信息 $M \in \mathbb{R}^{2d \times t}$ 与评论句子字符嵌入通过 Bi-LSTM 得到的隐表示 $H^c \in \mathbb{R}^{2d \times t}$ 进行点乘, 再和评论句子词嵌入通过 BiLSTM 得到的隐表示 $H^w \in \mathbb{R}^{2d \times t}$ 进行简单拼接, 得到最终的表征 $G \in \mathbb{R}^{2d \times t}$ 。

$$G = \{(M \cdot H^c) \oplus H^w\} \quad (12)$$

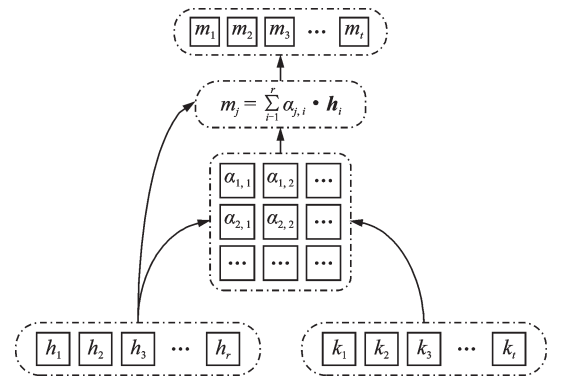


图2 交互注意力机制

Fig.2 Cross attention mechanism

式中“ $\cdot$ ”表示按位相乘。该表征既融合了关键词信息,又保留了原有的评论句字符级信息和词级信息在时序上的上下文依赖关系。

2.4 评价对象抽取层

将 $G=\{g_1, g_2, \cdots, g_t\}$ 通过一个线性层后,得到表示 G' ,其中 $G'_{i,j}$ 是序列中第 i 个字符的标签 j 的得分。设输入序列为 $x=\{x_1, x_2, \cdots, x_t\}$, 标签序列为 $y=\{y_1, y_2, \cdots, y_t\}$, 标签预测的分数为

$$\text{score}(x, y) = \sum_{i=0}^t Q_{y_i, y_{i+1}} + \sum_{i=1}^t G'_{i, y_i} \tag{13}$$

$$P(y|x) = \text{softmax}\{\text{score}(x, y)\} \tag{14}$$

式中: Q 为转移分数矩阵; $Q_{i,j}$ 表示从标签 i 转移到标签 j 的分数。对所有可能的标签序列的得分应用 softmax 函数,从而得到给定输入 x 的条件下标签序列 y 的概率 $P(y|x)$ 。本文使用负对数似然函数作为损失函数,最后通过维特比算法得到条件概率最大的输出序列。

$$L = - \sum_{x, y} \ln P(y|x) \tag{15}$$

$$\text{loss} = L + Z_i \tag{16}$$

式中 Z_i 表示第 i 个训练实例的惩罚项,见式(8)。

3 实验与分析

3.1 实验数据集

本文构建了两个涉案微博数据集,并对其评价对象进行人工标注,分别为“奔驰女司机维权案”与“重庆公交车坠江案”,具体如表3所示。

考虑到涉案微博评论数据的复杂性,本文对其进行了如下预处理:首先,去掉涉案微博评论数据集中重复的数据;其次,删除评论中出现的超链接广告,同时删除其中类似“@+用户名”的噪声数据;然后,如果评论中出现某些标点符号连续重复,比如“。。。”,“!!!”等,使用首位标点符号将其替换,同时删除评论中出现的表情符号;最后,为了确保涉案微博评论的完整性,删除其中字符数少于7的评论。

表3 涉案微博评价对象抽取数据集
Table 3 Data sets of cased-involved Microblog evaluation objects

数据集	训练集	验证集	测试集	总数量
	句子/评价对象	句子/评价对象	句子/评价对象	句子/评价对象
奔驰女司机维权案	1 005/797	120/123	125/130	1 250/1 050
重庆公交车坠江案	780/524	95/88	95/100	970/712

3.2 评价指标

本文的评价指标使用的是精确率(P)、召回率(R)、 F_1 值,计算公式分别为

$$P = \frac{TP}{TP + FP} \tag{17}$$

$$R = \frac{TP}{TP + FN} \tag{18}$$

$$F_1 = \frac{2 \times P \times R}{P + R} \tag{19}$$

式中:TP表示正例被正确判定为正例;FP表示负例被错误判定为正例;FN表示正例被错误判定为负例。

3.3 模型参数设置

本文实验所使用的预训练词向量^①是基于CTB 6.0(Chinese Treebank 6.0)语料库训练得到,字符嵌入^②是基于大规模标准分词后的中文语料库Gigaword训练得到,嵌入维度均为50。通过实验比较,选择关键词个数为20。

实验使用随机梯度下降算法(SGD)优化参数,dropout的大小设置为0.4,学习率设置为0.012, L_2 设置为 $1e-8$ 。

3.4 基准模型对比实验

为了验证本文模型的有效性,分别与CRF、LSTM-CRF、BiLSTM-CRF、BiLSTM-CNN-CRF、Seq2Seq4ATE和BERT-CRF六个基准模型进行了对比实验。基准模型介绍如下。

- (1)CRF^[8]:该方法是解决序列标注问题用得最多的方法之一,通过学习观察序列来预测标签序列。
- (2)LSTM-CRF^[20]:该方法也是序列标注问题中常用的方法,使用LSTM解决了远距离依赖问题。
- (3)BiLSTM-CRF^[20]:该模型使用BiLSTM从两个方向编码信息,来更好地捕获上下文信息,同时使用CRF向最终的预测标签添加约束。
- (4)BiLSTM-CNN-CRF^[21]:该模型结合了BiLSTM和CRF的优势,同时又融合了CNN抽取局部特征,进行评价对象抽取。
- (5)Seq2Seq4ATE^[13]:该模型使用门控单元控制编码器和解码器产生的隐状态,并通过位置感知注意关注目标词的相邻词。
- (6)BERT-CRF^[22]:该方法将评论句输入预训练BERT模型,得到的表示送入CRF,抽取评价对象词项。

为了保证比较的公平性,本文实验将上述模型的学习率、dropout、批次等参数设置为与本文模型一致,LSTM的隐层向量大小设置为100,CNN卷积核的尺寸设置为(2,3,4)。BERT-CRF实验中使用的BERT预训练语言模型为Google发布的BERT-Base(Chinese)模型。实验分别在两个数据集上进行,表4给出了对比实验的结果。通过表4可以看出,相较其他模型而言,基于传统机器学习的CRF模型的性能都是最低的,在两个数据集上的 F_1 值分别只有56.14%和45.81%,这是由于CRF模型需要定义大量的特征函数,根据自定义的语言特征模板进行评价对象抽取,并没有抽取相应的语义特征。与CRF模型相比,LSTM-CRF、BiL-

STM-CRF和BiLSTM-CNN-CRF模型利用LSTM对评论信息进行了抽取,因此性能获得了提升。其中,BiLSTM-CRF模型相较LSTM-CRF模型性能提升明显,这是由于BiLSTM是从前后两个方向编码信息,能够更好地捕获双向语义依赖关系,可以提取到某些很重要词的完整特征,而单

表4 基准模型对比实验结果							%
Table 4 Comparison of experimental results of baseline models							
模型	奔驰女司机维权案			重庆公交车坠江案			
	P	R	F_1	P	R	F_1	
CRF	57.14	55.17	56.14	44.44	47.27	45.81	
LSTM-CRF	63.57	54.67	58.78	49.11	50.00	49.54	
BiLSTM-CRF	70.99	64.14	67.39	62.11	50.43	55.66	
BiLSTM-CNN-CRF	70.50	67.59	69.01	57.55	55.45	56.48	
Seq2Seq4ATE	75.37	67.33	71.12	59.82	55.75	57.71	
BERT-CRF	76.74	68.28	72.26	62.79	55.89	59.12	
本文模型	78.79	69.33	73.75	69.07	57.26	62.61	

①<https://github.com/LeeSureman/Flat-Lattice-Transformer>
②<https://github.com/LeeSureman/Flat-Lattice-Transformer>

向的LSTM只能捕获到单向的词序列信息。融合了CNN模型之后, F_1 值又有所提升,说明CNN可以很好地捕获局部特征。与基于有CRF解码器的模型相比,Seq2Seq4ATE模型在性能上有进一步的提升,这是由于门控单元可以在解码当前词的标签时自动集成来自编码器和解码器隐藏状态的信息,同时位置感知模块可以迫使解码器在对标签进行解码时更加注意当前单词的邻近词。在基准模型中,基于预训练BERT的BERT-CRF模型的 P 、 R 、 F_1 值都是最高的,这是由于BERT包含了很多预训练语料中蕴含的外部知识和语义信息。在两个数据集上,本文模型的 P 、 R 、 F_1 值对比所有基准模型均有所提高,验证了本文模型抽取涉案微博评论的评价对象的有效性。

3.5 消融实验

为了验证本文模型中结构编码机制和案件关键词信息的有效性,针对“奔驰女司机维权案”数据集进行了消融实验。

第1个消融模型是“(-)案件关键词”,表示去掉模型中的关键词编码层和信息融合层,仅仅是对微博评论句子进行编码,将得到的多层次向量表示 K 送入CRF抽取评价对象词项。

第2个消融模型是“(-)结构编码”,表示去掉结构编码机制,仅仅将具有上下文语义特征的案件关键词的向量表示 L 与评论句子的多层次向量表示 K 做交互注意力,后续操作与主模型保持一致。两个消融模型的实验结果如表5所示。通过表5的实验结果可以看出,当没有融入案件关键词时,模型的 P 、 R 、 F_1 值均大幅下降,说明案件关键词的融入可以很好地指导模型学习涉案微博领域的特征,进而抽取评价对象词项。当没有使用结构编码机制时,模型的 F_1 值降低了1.34%, P 值降低了3.79%, R 值反而升高了0.67%,可以看出结构编码机制以牺牲一部分召回率换取了评价对象抽取精确率大的提升,说明结构编码机制可以有效帮助模型综合利用各个案件关键词的信息,对模型的指导作用更准确。

3.6 案件关键词数量的影响

本文针对两个数据集分别采用不同数量的案件关键词进行了实验,实验结果如图3所示。通过图3的实验结果可以看出,当案件关键词数量采用20和30时,性能相对较好。特别是当关键词数个数20时,模型在两个数据集上的 F_1 值都是最高的。说明当关键词数量过少时,其信息量不足,无法充分指导模型学习涉案微博领域的特征,而当关键词数量过多时,可能会引入噪声数据,让模型学习到错误的信息,导致模型性能下降。

3.7 案件关键词质量的影响

为了探究案件关键词质量对模型的影响,本文分别采用TextRank和TF-IDF两种关键词提取方法进行了实验。由于上述实验结果证明,抽取20个关键词融入模型的效果最好,所以从“奔驰女司机维权案”数据集的正文中利用两种方法分别抽取20个关键词,抽取结果如表6所示。通过表6可以看出,TextRank抽取到的关键词信息与正文中所提及的案件核心要素更相关,而TF-IDF会抽取到一些高频的噪声词,比如

表5 消融实验结果对比
Table 5 Comparison of ablation experiment results %

模型	P	R	F_1
(-)案件关键词	71.43	66.67	68.96
(-)结构编码	75.00	70.00	72.41
本文模型	78.79	69.33	73.75

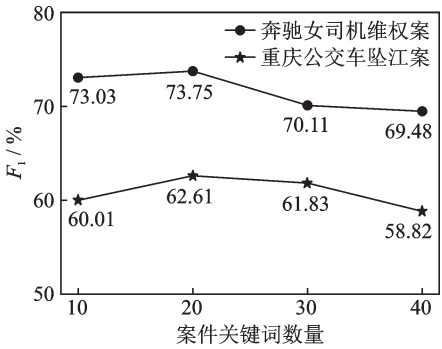


图3 设置不同关键词数量的 F_1 值对比

Fig.3 Comparison of F_1 values of different numbers of keywords

表 6 不同工具抽取到的案件关键词
Table 6 Case keywords extracted by different tools

案件	抽取工具	关键词
奔驰女司机 维权案	TextRank	消费者、监管部门、女士、奔驰、店方、媒体、当事人、记者、4s 店、执法人员、工商、车主、维权、利之星、销售、负责人、企业、品牌、经销商、法律法规
	TF-IDF	奔驰、视频、西安、女士、热议、高新区、监管部门、消费者、公正处理、利之星、4s 店、发动机、称有、营商、调查核实、奔驰车、店方、市场、当事人、记者

“热议、称有”等。

将表 6 得到的不同质量的关键词融入模型进行实验,实验结果如表 7 所示。表 7 的实验结果证明了使用 TextRank 抽取关键词的效果要优于 TF-IDF。原因可能是通过 TF-IDF 抽取到的关键词包含很多与评价对象无关的噪声词,这些词并不构成网友评论的评价对象,影响了模型的性能。

4 结束语

本文提出了一种基于关键词结构编码的涉案微博评价对象抽取方法。该方法通过结构编码机制综合利用微博正文的案件关键词信息,并通过交互注意力机制将其融入评论句子表示,来指导评价对象的抽取。在两个数据集上的实验结果证明了本文所提方法在涉案微博评价对象抽取方面的有效性。所提出的结构编码机制使模型能够更准确地抽取评价对象词项,且使用 TextRank 抽取一定数量的关键词融入模型能够得到最好的性能。未来的研究工作考虑融入其他涉案领域外部知识来改善评价对象抽取的结果。

参考文献:

[1] HU Mingqing, LIU Bing. Mining and summarizing customer reviews[C]//Proceedings of the 10th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. New York: ACM, 2004: 168-177.

[2] ZHUANG Li, JING Feng, ZHU Xiaoyan. Movie review mining and summarization[C]//Proceedings of the 15th ACM International Conference on Information and Knowledge Management. New York: ACM, 2006: 43-50.

[3] BLAIR-GOLDENSOHN S, HANNAN K, MCDONALD R, et al. Building a sentiment summarizer for local service reviews [C]//Proceedings of the Workshop on NLP in the Information Explosion Era.Stroudsburg, PA: ACL, 2008: 339-348.

[4] QIU Baojun, ZHAO Kang, MITRA P, et al. Get online support, feel better-sentiment analysis and dynamics in an online cancer survivor community[C]//Proceedings of the 3rd International Conference on Social Computing. Piscataway, NJ: IEEE, 2011: 274-281.

[5] KESSLER J S, NICOLOV N. Targeting sentiment expressions through supervised ranking of linguistic configurations[C]// Proceedings of the 3rd International AAAI Conference on Weblogs and Social Media. Menlo Park, California: AAAI, 2009: 90-97.

[6] JIN Wei, HUNG H H, SRIHARI R K. A novel lexicalized hmm-based learning framework for web opinion mining[C]// Proceedings of the 26th Annual International Conference on Machine Learning. New York: ACM, 2009: 465-472.

[7] JAKOB N, GUREVYCH I. Extracting opinion targets in a single and cross-domain setting with conditional random fields[C]// Proceedings of the 2010 Conference on Empirical Methods in Natural Language Processing. Stroudsburg, PA: ACL, 2010: 1035-1045.

[8] SHU Lei, XU Hu, LIU Bing. Lifelong learning crf for supervised aspect extraction[C]//Proceedings of the 55th Annual

表 7 不同质量关键词的实验结果对比

Table 7 Comparison of experimental results of different quality keywords %

抽取关键词方法	<i>P</i>	<i>R</i>	<i>F₁</i>
TF-IDF	75.74	68.67	72.02
TextRank	78.79	69.33	73.75

- Meeting of the Association for Computational Linguistics. Stroudsburg, PA: ACL, 2017: 148-154.
- [9] LI Yanzeng, LIU Tingwen, LI Diying, et al. Character-based BiLSTM-CRF incorporating POS and dictionaries for Chinese opinion target extraction[C]//Proceedings of the 10th Asian Conference on Machine Learning. Vienne, Austria: PMLR, 2018: 518-533.
- [10] LI Xin, LAM W. Deep multi-task learning for aspect term extraction with memory interaction[C]//Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing. Copenhagen: EMNLP, 2017: 2886-2892.
- [11] LI Xin, BING Lidong, LI Piji, et al. Aspect term extraction with history attention and selective transformation[C]//Proceedings of the 27th International Joint Conference on Artificial Intelligence. Menlo Park, CA: AAAI, 2018: 4194-4200.
- [12] WANG Wenya, PAN S J, DAHLMEIER D, et al. Coupled multi-layer attentions for co-extraction of aspect and opinion terms [C]//Proceedings of the 31st AAAI Conference on Artificial Intelligence. Menlo Park, CA: AAAI, 2017: 3316-3322.
- [13] MA Dehong, LI Sujian, WU Fangzhao, et al. Exploring sequence-to-sequence learning in aspect term extraction[C]//Proceedings of the 57th Conference of the Association for Computational Linguistics. Stroudsburg, PA: ACL, 2019: 3538-3547.
- [14] HE Ruidan, LEE W, NG H, et al. An unsupervised neural attention model for aspect extraction[C]//Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics. Stroudsburg, PA: ACL, 2017: 388-397.
- [15] LUO Ling, AO Xiang, SONG Yan, et al. Unsupervised neural aspect extraction with sememes[C]//Proceedings of the 28th International Joint Conference on Artificial Intelligence. Menlo Park, CA: AAAI, 2019: 5123-5129.
- [16] TULKENS S, CRANENBURGH A V. Embarrassingly simple unsupervised aspect extraction[C]//Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics. Stroudsburg, PA: ACL, 2020: 3182-3187.
- [17] LI Kun, CHEN Chengbo, QUAN Xiaojun, et al. Conditional augmentation for aspect term extraction via masked sequence-to-sequence generation[C]//Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics. Stroudsburg, PA: ACL, 2020: 7056-7066.
- [18] LIN Zhouhan, FENG Minwei, DOS SANTOS C N, et al. A structured self-attentive sentence embedding[EB/OL]. (2017-03-03). <https://arxiv.org/abs/1703.03130>.
- [19] SRIVASTAVA R K, GREFF K, SCHMIDHUBER J. Highway networks[EB/OL]. (2015-05-03). <https://arxiv.org/abs/1505.00387>.
- [20] HUANG Zhiheng, XU Wei, YU Kai. Bidirectional LSTM-CRF models for sequence tagging[EB/OL]. (2015-08-01). <https://arxiv.org/abs/1508.01991>.
- [21] MA Xuezhe, HOVY E. End-to-end sequence labeling via Bi-directional LSTM-CNNs-CRF[C]//Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics. Stroudsburg, PA: ACL, 2016: 1064-1074.
- [22] DEVLIN J, CHANG Mingwei, 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 Computational Linguistics. Stroudsburg, PA: ACL, 2019: 4171-4186.

作者简介:



王静赞(1996-),女,硕士研究生,研究方向:自然语言处理、情感分析,E-mail: 18803970927@163.com。



余正涛(1970-),男,教授,研究方向:自然语言处理、信息检索、机器翻译、机器学习。



相艳(1979-),通信作者,女,副教授,研究方向:文本挖掘和情感分类,E-mail: sharonxiang@126.com。



陈龙(1994-),男,硕士研究生,研究方向:自然语言处理、文本挖掘。

(编辑:夏道家)