

DOI:10.11918/201907004

# 融合表情符号与短文本的微博多维情感分类

赵晓芳,金志刚

(天津大学 电气自动化与信息工程学院,天津 300072)

**摘 要:** 表情符号已成为网络语言重要组成部分,是分析社交媒体情感的主要特征之一. 目前分析社交媒体情感符号的方法多针对 Emoji,对颜文字的情感倾向没有相应分析. 为获取中文媒体的多维度情感并分析热点话题的群体情感走向,本文以微博为例提出一种新的融合表情符号与短文本的多维情感分类方法. 在该框架中,采用深度学习模型分析文本与 Emoji 组合部分、颜文字部分,分别计算两部分的 7 种情感强度,挖掘各部分与情感标签的深层次关联,并设计计算模型来反映语句包含的多维情感属性,实现对语句多维情感强度的检测. 实验选择 NLPCC2014 数据集和爬取的带有颜文字的微博数据集进行验证,实验证明当文本与 Emoji 组合、颜文字占比分别为 0.6 和 0.4 时情感分类效果最好,且含颜文字的语句情感分类性能指标始终高于不含颜文字的语句,这表明融合表情符号和短文本的形式有效提高了情感检测精度. 该方法为研究群体情感趋势提供了更细粒度的分析,为中文社交媒体的情感分析提供了新思路.

**关键词:** 情感分类;Emoji;颜文字;深度学习

**中图分类号:** TP391      **文献标志码:** A      **文章编号:** 0367-6234(2020)05-0113-08

## Multi-dimensional sentiment classification of microblog based on Emoticons and short texts

ZHAO Xiaofang, JIN Zhigang

(School of Electrical and Information Engineering, Tianjin University, Tianjin 300072, China)

**Abstract:** Emoticons have become an important component of network language and is one of the main characteristics of the analysis of social media sentiment. The current social media sentiment analysis methods most focus on Emoji, while there is no study on the sentiment trend of kinesics. In order to obtain the multi-dimensional sentiment polarity of Chinese social media and analyze the group sentiment trend on hot topics, this paper proposes a new multi-dimensional sentiment classification method based on deep learning, which combines Emoticons with short texts. In this framework, the text and Emoji combination and the kinesics in microblog sentences were analyzed using deep learning model, and seven sentiment intensities of the two parts were obtained to explore the correlation between each part and sentiment labels. Then, a computational model was designed to reflect the multi-dimensional sentiment polarity contained in microblog sentences, which can realize the detection of the multi-dimensional sentiment intensity of sentences. The experiment utilized the NLPCC2014 dataset and the crawled microblog dataset containing kinesics for verification. Results show that when the proportion of the text and Emoji combination and the kinesics were 0.6 and 0.4, the effect of sentiment classification was the best. The sentiment classification performance indicator of the sentences containing kinesics was always higher than that without kinesics, which indicates that the combination of Emoticon and short texts can effectively improve the accuracy of microblog sentiment detection. The experiment provides a more fine-grained analysis for group sentiment trend and a new idea for Chinese social media sentiment analysis.

**Keywords:** sentiment classification; Emoji; kinesics; deep learning

2019 年北京大学心理与认知科学学院发表《95 后手机使用心理与行为白皮书》<sup>[1]</sup>,调研报告显示,我国 95 后网民占比高达 27.9%,越来越多的 95 后通过社交平台来表达自己的观点和意见,且 95 后使

用表情包的意愿度为 69.8%,因此在社交网络评论中会出现大量表情符号使表达的观点更真实更丰富. 微博平台为用户提供了大量的默认表情字符,如 😊,😂 等,帮助用户更生动地表达当下的感受和体会. 此外,字符形态的表情“颜文字”是另一种常用的表情符号,由于多语言的文化包容性和字符多样性,已发展成为影响全球的文化符号.

表情符号的使用可以揭示隐藏在文本之下的情

收稿日期: 2019-07-01  
基金项目: 国家自然科学基金(71502125)  
作者简介: 赵晓芳(1990—),女,博士研究生;  
金志刚(1972—),男,教授,博士生导师  
通信作者: 金志刚, zgjin@tju.edu.cn

感,例如,“体测一千米,都没有喘不过气来”表示快乐和兴奋;“我不知道把钢笔放在了哪里 QAQ”为不安的情绪;“连续 42 小时的工作,已经忘记了时间的存在”表示兴奋和收获的满足.在这些句子中,虽然文字并没有明确地传达出情绪,但结尾处的表情符号可以有效识别个体的情感状态.

2019 年 3 月,新浪微博数据中心发布最新《2018 微博用户发展报告》<sup>[2]</sup>,报告显示 2018 年第 4 季度微博月活跃用户 4.62 亿,其中 95 后用户高达 41%.由于微博有着迅捷性、蔓延性、平等性与组织性等 4 大特点,其热点话题随时都会引发万级的转发和评论,例如从“长生生物等疫苗造假事件”到“翟天临被指论文抄袭事件”再到最近的“996 工作制”事件,无一不是在微博上迅速发酵,并最终对现实社会产生影响,而且这种线上影响线下的趋势越来越明显.为了更好地利用微博,产生有益社会价值,消除潜在危害,本文提出微博多维度情感分析,这有利于分析群体情感倾向,提高舆情分析、引导的准确率,在网络言论尚未形成舆论前,及早预测可能导致用户负面情绪的舆论或报道,实现对舆情的早期介入和有效引导.

目前,针对中文媒体中情感符号进行分析的方法多针对 Emoji,对于颜文字情感倾向并没有相关分析.本文在文本基础上,结合了 Emoji、颜文字情感特征,提出了一种融合表情符号与短文本的多维情感分类方法(EmoT):收集 95 后在微博发表的评论,通过现有的日式颜文字资源,构建中文颜文字词典,提取颜文字中的结构特征、类别特征以及运动学特征,并利用多层感知机(Multi-Layer Perception, MLP)对颜文字进行情感检测,得到颜文字的 7 种情感强度值;采用基于注意力机制的 CNN 和 LSTM 对文本和 Emoji 进行编码,将得到的融合特征作为 MLP 的输入,进一步挖掘文本和 Emoji 组合部分与情感标签的深层次关联,计算文本和 Emoji 组合部分的 7 种情感强度值;结合以上两种情感倾向设计计算模型,通过相关实验,验证了构建的融合表情符号与短文本的多维情感分类方法可以进一步提高情感分析性能,补充了中文媒体评论中颜文字情感研究的空白,且该方法获取的语句 7 种情感强度,为群体情感走向提供了更细粒度分析.

## 1 相关工作

微博语句的情感分析,不仅需要对本微博文本进行分析,还需要考虑用户使用的表情符号,有相关研究将表情符号作为特征进行情感分析,这些方法改

善了社交媒体的情感分析结果,例如 Jonathon<sup>[3]</sup>提出使用朴素贝叶斯和支持向量机分类器代替话题敏感词进行情感分析时,可以使用表情符号来减少话题依赖性. Yang<sup>[4]</sup>等提出将 Emoji 这种表情符号作为自动注释工具,使用支持向量机和条件随机场等算法将句子标记为 4 种情感类别,可以减少手动注释节省时间和人力.

以上研究证实了表情符号对情感分析的积极影响.然而,对中文社交媒体的分析仅仅局限于 Emoji,并没有分析 Emoji 和颜文字两种表情符号对情感表达的影响.例如 Song 等<sup>[5]</sup>提出利用情感种子词和情感符号来构建情感词典,以便更准确地捕捉候选情感词的细粒度情感倾向;Jiang 等<sup>[6]</sup>利用表情符号的词向量构建表情符号向量空间模型,并将文本中所有的词映射到向量空间中,采用支持向量模型实现对文本的情感分类;何等<sup>[7]</sup>为常见的表情符号构建情感空间的特征表示矩阵,通过将文本词向量矩阵与表情符号的特征表示矩阵进行乘积运算,实现词义到情感空间的映射,从而构建文本的情感表示矩阵;张等<sup>[8]</sup>提出基于双重注意力模型的情感分析方法,分别对文本和情感符号进行编码.此外,相关学者对日式颜文字进行了情感分类,如 Utsu 等<sup>[9]</sup>通过对日式颜文字中包含的眼睛、嘴巴等符号的出现概率进而对颜文字进行情感极性判别;Yu 等<sup>[10]</sup>通过构建颜文字词典对日语旅游网站评论进行情感分类;Yu 等<sup>[11]</sup>提出了 AZEmo 系统,实现了对社交媒体和电子商务网站中颜文字的提取和分类.

本文针对包含 Emoji 和颜文字的语句,将其分为两部分进行多维情感分类.第一部分是对文本和 Emoji 组合部分进行情感倾向分析,为了提高特征表达能力,进一步挖掘文本和 Emoji 组合部分与情感的深层次关联,采用基于注意力结构的 CNN、LSTM 分别提取组合部分的语义特征,将得到的两种语义特征进行融合后,通过 MLP 进一步提取高级语义表示,最后使用 softmax 分类层得到文本与 Emoji 组合部分的多维情感强度值;第二部分通过构建的中文颜文字词典,提取颜文字结构、情感类别、运动学特征,采用 MLP 对颜文字的情感倾向进行分析,得到颜文字的 7 种情感强度值.最后基于上述两部分的情感倾向分析,设计情感计算模型得到语句的多维情感强度值.

## 2 融合表情符号与短文本的微博多维情感分类

新浪微博等社交媒体中的 Emoji 和颜文字是目前常用的两种表情符号,人们借助情感符号表达更

加微妙的情感变化,如加强、澄清或强调、幽默等.Emoji、颜文字是微博文本情感倾向的重要特征,相比于文字,表情符号对情感的语义区分能力更高.本文在对文本、Emoji 组合部分情感分析的基础上,结合了颜文字的情感倾向,为包含颜文字、Emoji 的社交媒体评论的情感分析提供了模型基础,如图 1 所示为情感分类模型框架.

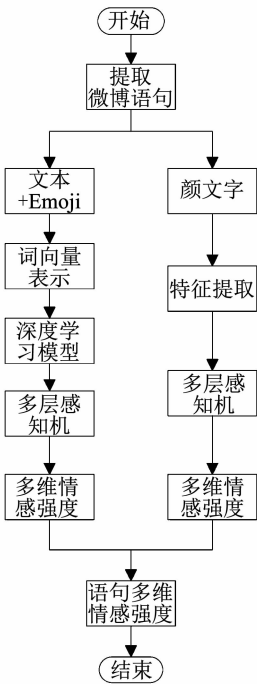


图 1 情感分类模型 (EmotT)

Fig. 1 Sentiment classification model (EmotT)

2.1 颜文字情感计算

颜文字的情感计算主要包括 3 部分:中文颜文字词典构建、颜文字提取以及颜文字情感倾向值计算.

2.1.1 颜文字词典构建

目前,并没有大规模的中文颜文字词典,因此本文在 Ptaszynski 等<sup>[12]</sup>构建日语颜文字词典基础上构建中文颜文字词典.在词典中,每种情感标签均包含一系列颜文字,且每个颜文字只属于一种情感标签.中文颜文字词典构建规则分以下两步:

明确颜文字情感标签:在构建中文颜文字词典时忽略无法推断出是否具有明显情感倾向的颜文字,只收集带有明确情感含义标签的颜文字.

截断处理:由于日式颜文字包含的不仅有眼睛、嘴巴等符号,还有一些日语来帮助解释颜文字的含义,如颜文字“ヾ(@へへ@)ノおはよう”中“おはよう”是早上好的意思,用来帮助解释颜文字的含义,在构建中文颜文字词典时,对其中的日语进行了截断处理,例如:日式颜文字“ヾ(@へへ@)ノおはよう”在构建的中文颜文字词典时最终变为“ヾ(@へへ@)ノ”.

2.1.2 颜文字提取

Birdwhistell<sup>[13]</sup>最先提出了人体运动学理论,该理论确定了颜文字中的字符所对应的身体或面部表情.本文采用 Michal 等<sup>[10]</sup>构建的人体运动学模型,将每个颜文字表示为 9 个部分: {S<sub>1</sub>} {B<sub>1</sub>} {S<sub>2</sub>} {E<sub>L</sub>} {M} {E<sub>R</sub>} {S<sub>3</sub>} {B<sub>2</sub>} {S<sub>4</sub>},各个部分与颜文字的对应关系如表 1.

表 1 颜文字组成  
Tab. 1 Components of kinesics

| 颜文字      | S <sub>1</sub> | B <sub>1</sub> | S <sub>2</sub> | E <sub>L</sub> | M | E <sub>R</sub> | S <sub>3</sub> | B <sub>2</sub> | S <sub>4</sub> |
|----------|----------------|----------------|----------------|----------------|---|----------------|----------------|----------------|----------------|
| T~T      | 空              | 空              | 空              | T              | ^ | T              | 空              | 空              | 空              |
| (-▽-;)   | 空              | (              | 空              | —              | ▽ | —              | ;              | )              | 空              |
| ヾ(@へへ@)ノ | ヾ              | (              | @              | へ              | — | へ              | @              | )              | ノ              |

{B<sub>1</sub>} {B<sub>2</sub>} 分别为颜文字的左右边界; {E<sub>L</sub>} {M} {E<sub>R</sub>} 与人体运动学中的左眼睛、嘴巴、右眼睛对应,是颜文字的核心部分并记为 {S<sub>5</sub>}; {S<sub>1</sub>} {S<sub>2</sub>} {S<sub>3</sub>} {S<sub>4</sub>} 为颜文字中与人体运动学理论对应的其他部分,如胳膊、汗水等,称为补充分量.并不是每个颜文字都严格地遵循这 9 个部分,没有的部分用“空”表示.中文文本中提取颜文字步骤如图 2.

2.1.3 颜文字情感倾向值计算

本文采用 Xu 等<sup>[14]</sup>提出的 7 大情感类别,将颜文字的情感分为乐、好、哀、恶、怒、惧和惊.此外,分别将颜文字的核心部分 {S<sub>5</sub>} 和 4 个补充分量 {S<sub>1</sub>} {S<sub>2</sub>} {S<sub>3</sub>} {S<sub>4</sub>} 映射到 7 种情感类别中,每个分量出

现在 7 种情感类别中概率计算为

$$P_{S_i}^c = N_{S_i}^c / N_{S_i}, (i = 1, 2, \dots, 5).$$
 (1)

式中: N<sub>S<sub>i</sub></sub> 为分量 {S<sub>i</sub>} 在颜文字中出现的总次数, C 为乐、好、哀、恶、怒、惧和惊 7 种情感标签集合, N<sub>S<sub>i</sub></sub><sup>c</sup> 为 {S<sub>i</sub>} 在标签 c 中出现的次数.

在文本挖掘和情感分析研究中,提取结构、句法和语义特征被广泛采用.然而,由于颜文字的性质不同于语言句子,一些特征(如句法特征)不适用于颜文字.因此,采用 Yu 等<sup>[11]</sup>提出的颜文字特征表示方法,从颜文字中提取 4 种结构特征、7 种分类特征和 35 种人体动力学特征.提取的特征以及解释如表 2.

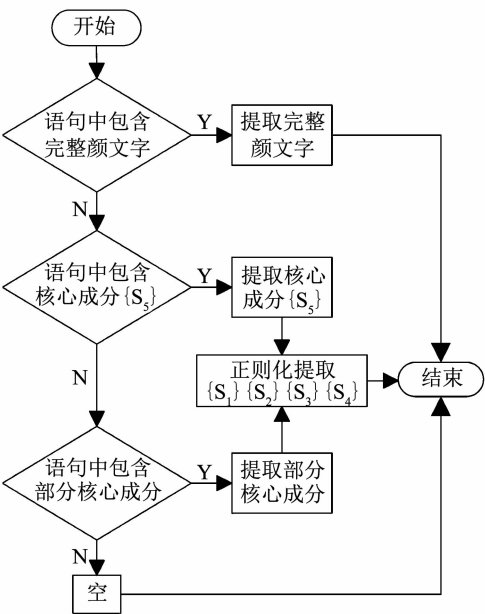


图 2 中文文本中提取颜文字框架

Fig. 2 Model of extracting kinesics from texts

表 2 颜文字特征

Tab. 2 Features of kinesics

| 特征类型 | 特征                                                                                                                                                                    | 解释                       |
|------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------|
| 结构特征 | Length                                                                                                                                                                | 颜文字长度                    |
|      | P <sub>ASCII</sub>                                                                                                                                                    | ASCII 字符比例               |
|      | N <sub>charType</sub>                                                                                                                                                 | 不同类型字符的数量                |
| 类别   | N <sub>maxRecurringCharCount</sub>                                                                                                                                    | 出现次数最多的字符数量              |
|      | P <sup>h</sup> , P <sup>s</sup> , P <sup>f</sup> ,                                                                                                                    | 所有字符情感类别概率和              |
| 特征   | P <sup>a</sup> , P <sup>d</sup> , P <sup>u</sup> , P <sup>l</sup>                                                                                                     |                          |
| 运动学  | P <sup>h</sup> <sub>S<sub>i</sub></sub> , P <sup>s</sup> <sub>S<sub>i</sub></sub> , P <sup>f</sup> <sub>S<sub>i</sub></sub> ,                                         | {S <sub>i</sub> } 情感类别概率 |
| 特征   | P <sup>a</sup> <sub>S<sub>i</sub></sub> , P <sup>d</sup> <sub>S<sub>i</sub></sub> , P <sup>u</sup> <sub>S<sub>i</sub></sub> , P <sup>l</sup> <sub>S<sub>i</sub></sub> |                          |

结构特征考虑了颜文字的结构信息. 该类型包括颜文字的总长度、ASCII 字符的比例、不同类型字符的数量以及颜文字中出现频率最高的字符的数量.

情感特征借助 Yamada 等<sup>[15]</sup>提出的方法,通过提取颜文字中每个字符在特定情感类别中的出现次数除以该字符在所有颜文字中的出现总次数,得出该字符表达特定情感的概率.

通过从表情中提取人体运动学特征,计算核心和其他成分的情感类别概率. 颜文字的情感是由人体运动学成分构成的,这意味着运动学成分的情感概率在颜文字的情感类别中占很大比例. 如果两个人体运动学成分在两个情感类别中具有相似的类别频率,则认为它们在构建表情符号方面具有相似的功能. 若颜文字中不包含某些人体运动学成分时,其对应的情感概率为 0. 通过上述分析,特征向量中共

有 35 个人体运动学特征值(有些为空).

情感预测质量的评价与具体的预测模型有关,本文采用不含隐藏层的 MLP 分析颜文字的情感倾向. 对 MLP 的输出向量采用 Rectifier 函数进行非线性变换后,得到情感标签的得分向量为

$$S_{\text{score}}(K) = g(\mathbf{W}_K \mathbf{x} + \mathbf{b}_K). \tag{2}$$

式中:  $S_{\text{score}}(K) \in \mathbf{R}_{|C|}$  为颜文字  $K$  的情感得分向量,  $C$  为 7 种情感标签集合,  $g$  取 RELU 函数,  $\mathbf{W}_K$ 、 $\mathbf{b}_K$  分别为 MLP 的参数矩阵和偏置,情感得分向量作为 softmax 的输入,得出颜文字的多情感分类概率为

$$P_K(c/K) = \frac{\exp(S_{\text{score}}(K)_n)}{\sum_{n=1}^C \exp(S_{\text{score}}(K)_n)}. \tag{3}$$

得到多维情感强度为:  $Q_K = P_K(c/K)$ .

2.2 基于注意力的文本和 Emoji 情感计算

本文采用 Jiang<sup>[6]</sup>提出的利用表情符号构建表情符号向量空间模型,并将文本中的所有词映射到向量空间中. 文本词语及 Emoji 词向量的获取可以看作是一个查词典的过程. 词典  $\mathbf{R}^{d \times N}$  通过大规模语料采用词向量训练模型学习得到. 对于一个文本序列  $T = \{t_1, t_2 \cdots t_n\}$ , 将词语的词向量拼接起来,就得到整个文本序列的词向量表示,拼接方式如式(4)所示:

$$\mathbf{R}_T = \mathbf{r}_1 \oplus \mathbf{r}_2 \oplus \cdots \mathbf{r}_{|N|}. \tag{4}$$

式中:  $\mathbf{r}_i \in \mathbf{R}^{d \times N}$  为  $t_i$  对应于词典中的元素,  $\oplus$  为行向量拼接操作,  $d$  为词向量的维数,  $|N|$  为词典中词语的个数. 对 Emoji 集合中的词向量同样采用上式的方式进行拼接,  $\mathbf{R}_T$  和  $\mathbf{R}_E$  的维数分别为词语的数目和 Emoji 的数目.

文本的词向量矩阵  $\mathbf{R}_T$  以及 Emoji 的词向量矩阵  $\mathbf{R}_E$  作为深度学习模型的输入. 为了进一步提高对文本和 Emoji 建模能力,首先采用基于注意力机制的 CNN 和 RNN 分别提取局部以及上下文相关特征,然后对提取的特征进行融合,增强模型表达能力的同时获得更多相关的语义特征. 如图 3 为基于注意力的文本和 Emoji 建模框架.

基于 CNN 与 RNN 模型的语义获取与融合

$\mathbf{R}_T$  和  $\mathbf{R}_E$  作为 CNN 模型的输入,卷积层用大小为  $h \times d$  的滤波器对文本矩阵执行卷积操作,提取局部特征:

$$c_i = f(\mathbf{F} \cdot \mathbf{R} + \mathbf{b}). \tag{5}$$

式中:  $\mathbf{F}$  代表宽度为  $h$  的滤波器,  $\mathbf{b}$  为偏置量;  $f$  为通过 RELU 进行非线性操作的函数;  $c_i$  为通过卷积操作得到的局部特征. 随着滤波器依靠为 1 的步长从上往下进行滑动,采用 VALID 方式进行 padding 操作,获得与原输入相同长度的特征向量  $\mathbf{C}_i$ .

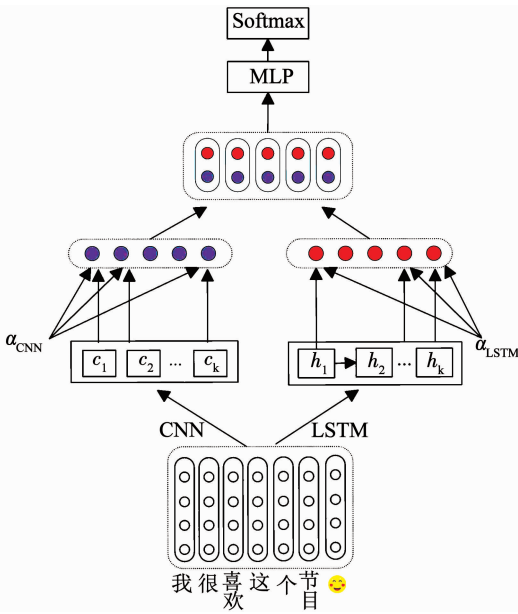


图3 基于注意力的文本和Emoji建模框架

Fig.3 Model of attention-based text and Emoji

对于输入序列  $x = (x_1, x_2, \dots, x_i)$ , 一个标准的 RNN 模型计算隐藏和输出序列分别为

$$h_t = \sigma(W_{sh}x_t + W_{hh}h_{t-1} + b_h), \quad (6)$$

$$y_t = W_{hy}h_t + b_y. \quad (7)$$

式中:  $W$  项为权重,  $b$  为偏置,  $\sigma$  取 sigmoid 函数. 然而, 由于 RNN 梯度容易消失、爆炸等缺点, RNN 的变体长短时记忆 (LSTM) 被越来越广泛的应用. LSTM 模型可看作一种权值共享的深度学习神经网络, 利用门信号的线性循环连接来解决时间层面的梯度弥散问题, 进而学习长期的依赖关系. 本文使用 LSTM 模型提取语句的语义信息表示为  $H_t$ . 为捕获更直接的语义依赖关系, 采用 Bahdanau<sup>[16]</sup> 提出的注意力模型, 将两种模型的输出分别作为注意力模型的输入, 得到的文本表示向量为

$$v = \sum_{i=1}^k \alpha_i f_i. \quad (8)$$

式中  $f_i$  为  $C_i$  或  $H_i$ ; 输入状态的权重为

$$\alpha_i = \frac{\exp(e_i)}{\sum_{i=1}^k \exp(e_i)}. \quad (9)$$

式中:  $e_i = \text{Tanh}(W h_i + b)$ ,  $W$  为模型权重,  $b$  为模型偏置. 因此得到的基于注意力机制的语义信息分别表示为  $v_i^C$  和  $v_i^L$ . 在得到基于注意力的语义信息后, 将两种信息融合为:  $v = [[v_1^C, v_1^L], [v_2^C, v_2^L], \dots, [v_k^C, v_k^L]]$ .

为了得出文本与 Emoji 融合特征的多分类情感概率, 与处理颜文字的过程类似, 将融合特征作为 MLP 的输入, 并对 MLP 的输出向量采用 Rectifier 函数进行非线性变换, 得到情感标签的得分向量:

$$S_{\text{score}}(S) = g(W_{\text{TE}}v + b_{\text{TE}}). \quad (10)$$

式中:  $S_{\text{score}}(S) \in R_{|C|}$  为文本与 Emoji 融合部分的情感分类向量,  $C$  为 7 种情感标签集合,  $g$  取 RELU 函数,  $W_{\text{TE}}$ 、 $b_{\text{TE}}$  分别为 MLP 的参数矩阵和偏置, 情感得分向量作为 softmax 的输入, 得出文本与 Emoji 的多情感分类概率为

$$P_{\text{TE}}(c/S) = \frac{\exp(S_{\text{score}}(S)_i)}{\sum_{i=1}^L \exp(S_{\text{score}}(S)_i)}. \quad (11)$$

选择 softmax 层输出的概率作为多维情感强度值  $Q_{\text{TE}} = P_{\text{TE}}(c/S)$ .

### 2.3 微博语句情感计算

对于情感倾向要综合考虑表情符号和文本两部分, 通过对上述两部分的情感倾向值加权处理后即可得到微博的多维情感概率:  $Q = (1 - \lambda) Q_{\text{TE}} + \lambda Q_K$ .  $Q_{\text{TE}}$  为微博文本与 Emoji 的多维情感概率,  $Q_K$  为颜文字的多维情感概率,  $\lambda$  为正的可变参数, 表示情感倾向所占的比重, 且  $\lambda \in (0, 1)$ .

## 3 实验与分析

本文针对 95 后在社交网络中发表的评论进行情感倾向性分析, 考虑到 95 后在社交媒体的评论中喜欢使用颜文字的特点, 提出的情感分类模型考虑了文本中包含的 Emoji 及颜文字, 由于目前并未发现合适的公开数据集, 因此自行爬取 2017 年—2018 年约 40 000 000 条微博评论构建实验数据集, 并将其分为带有颜文字的与不带有颜文字的两类. 由于 Emoji 和颜文字数量较多, 通过对表情符号进行情感标注和降序排列, 最终各选择前 100 的表情符号进行实验. 实验进行了两部分的对比, 第一组是滤除颜文字, 采用提出的基于注意力机制的方法与选取的模型进行对比, 验证提出模型对文本和 Emoji 组合部分建模的有效性; 第二组是对含有 Emoji、颜文字的语句进行多情感分类, 验证颜文字对情感分类性能的影响.

### 3.1 第一组实验: 无颜文字

实验数据采用 NLPCC 的 2014 年中文微博评测任务的公共数据集 NLPCC2014. 数据集分为测试集和训练集两部分, 标注的情感标签主要分为 7 类, 如表 3 所示:

为了学习文本词语和 Emoji 的词向量, 本文将爬取的不含有颜文字的语句经过 NLPIR 分词、去噪、繁简转换、url 替换以及删除长度较短的无意义的微博之后, 构建了 Word2Vec 词向量训练语料库, 包含微博 31 394 583 条, 词语数 1 034 869 204 个, 选择 skip-gram 模型进行训练, 得到词向量空间包含

649 302 个词,词向量维数为 200 维. 实验中 CNN、LSTM 的参数如表 4 所示.

表 3 NLPCC2014 标注数据集  
Tab.3 Dataset of NLPCC2014

| 情感标签         | 训练集   | 测试集   | 总计    |
|--------------|-------|-------|-------|
| 乐(happiness) | 1 459 | 441   | 1 900 |
| 好(like)      | 2 204 | 1 042 | 3 246 |
| 哀(sadness)   | 1 173 | 189   | 1 362 |
| 惧(fear)      | 148   | 46    | 194   |
| 怒(anger)     | 669   | 128   | 797   |
| 恶(disgust)   | 1 392 | 389   | 1 781 |
| 惊(surprise)  | 362   | 1 662 | 524   |

表 4 模型参数  
Tab.4 Parameters of CNN and LSTM

| 模型       | 参数               | 说明        | 取值       |
|----------|------------------|-----------|----------|
| CNN      | d                | 词向量维度/维   | 200      |
| CNN      | d <sub>win</sub> | 卷积算子宽度    | 4        |
| CNN      | H                | 卷积算子数/个   | 200      |
| LSTM     | Lstm_units       | 输出维度/维    | 200      |
| CNN/LSTM | Optimizer        | 模型优化器     | Adadelat |
| CNN/LSTM | batch_size       | batch 样本数 | 32       |

实验选取的对比模型如下:  
MNB 模型(multinomial naive bayes, MNB): 作为机器学习的代表,在很多情感分类任务中都取得了优秀的效果.

ESM 模型:Jiang 等<sup>[6]</sup>针对 Emoji 提出的表情符号空间模型,完成了词语到情感空间的映射,本文选择对映射后的词语求最大、最小和求和的方法作为对比模型.

EMCNN 模型:何等<sup>[7]</sup>提出的为 Emoji 构建语义特征表示矩阵,通过矩阵乘积运算增强微博文本语义,采用多通道卷积神经网络对特征学习实现情感分类.

DAM 模型:张等<sup>[8]</sup>提出的基于双重注意力机制,分别对文本和情感符号进行编码,在构建语义表示后进行情感分类.

通过实验,对 NLPCC2014 评测结果如表 5 所示:

表 5 NLPCC2014 评测结果  
Tab.5 Results for NLPCC2014

| 模型          | F1            |               |
|-------------|---------------|---------------|
|             | MicroF1       | MacroF1       |
| MNB         | 0. 359        | 0. 278        |
| ESM         | 0. 442        | 0. 378        |
| EMCNN       | 0. 472        | 0. 394        |
| DAM         | 0. 494        | 0. 414        |
| DA-LSTM/CNN | <b>0. 518</b> | <b>0. 432</b> |

从上表看出,本文方法得到的 MicroF1 和 MacroF1 两个评价指标较其他模型分别平均提高 7. 6% 和 6. 6%,明显优于其他模型,这主要是因为采用基于注意力机制的 CNN 和 RNN 分别提取局部以及上下文相关特征,并对提取的特征进行融合,进一步提高了对文本和 Emoji 建模能力. 其次,为了进一步验证该模型可以在增强模型表达能力的同时获得更多相关的语义特征,选择了两个实例来可视化注意力结果. 可视化结果如图 4 和图 5 所示. 其中颜色由浅到深意味着注意力权重由高到低.

如图 4 所示,“这是一部令人兴奋的关于运动的动作电影,还有一个我喜欢的故事🤔”,看出本文提出的模型为每个特征提取层分配了不同的注意力权重. 在基于 LSTM 的特征提取层中,兴奋、[乐]和喜欢的权重更高,在基于 CNN 的特征提取层中,故事、兴奋和[乐]的权重更高,而且在两种特征提取中,Emoji 都具有较高的注意力权重,CNN 和 LSTM 模型都可以很好地捕捉语句中的表情符号,最后该模型将它们结合起来确定情感极性. 此外,句子中的一些噪音如“的”、“令人”等对语句情感极性判断没有贡献,因此其注意力值很小.

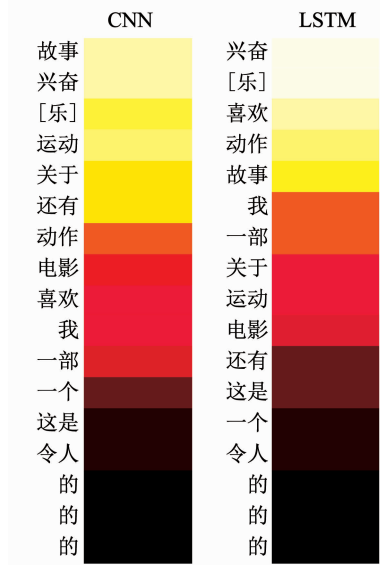


图 4 注意力可视化结果(“这是一部令人兴奋的关于运动的动作电影,还有一个我喜欢的故事[乐]”.)

Fig. 4 Visualization of attention (“This is an exciting action movie about sports, and there is a story I like [laugh].”)

在图 5 中,“几乎在每个关键时刻影片都破坏了剧本的亮点🤔”. 破坏不是形容词,却表达了隐含的消极极性. 同时可以发现基于 LSTM 的特征提取中,破坏和[怒]权重较高,基于 CNN 的特征提取中,关键时刻权重较高,而且在 LSTM 中 Emoji 的注意力权重高于 CNN 中的注意力权重,两者对表情符

号的注意力程度不同,最后模型融合了以上部分确定情感极性。

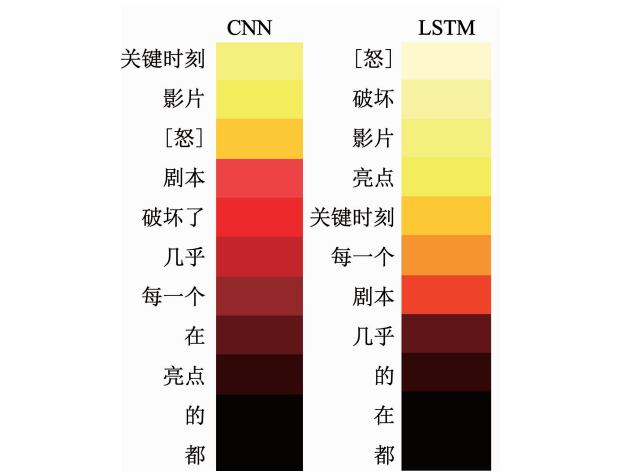


图5 注意力可视化结果(“几乎在每个关键时刻影片都破坏了剧本的亮点[怒]”)

Fig. 5 Visualization of attention (“The highlights of the script are destroyed by the movie at almost every important moment[anger].”)

上述两个可视化结果都反映了提出模型对语义表达能力的增强,即本文提出的基于注意力机制的深度学习模型,通过特征提取和特征组合两步将局部信息和上下文信息组合在一起,很好地捕捉文本词语、Emoji 和句子结构的搭配,以识别情感倾向。

### 3.2 第二组实验:包含颜文字

为验证提出的融合表情符号和短文本的多维情感分类效果,选取爬取的带有颜文字的微博语句 8 000 条和不带有颜文字的微博语句 2 000 条,分别验证不同的  $\lambda$  取值对情感分类准确率的影响,以及加入颜文字对最终情感分类准确率的影响。

#### 3.2.1 $\lambda$ 取值

对于微博的情感倾向,要综合考虑文本、Emoji 组合部分以及颜文字的情感,通过对这两部分的情感倾向值加权处理后即得到微博情感倾向值。为了判断  $\lambda$  的取值,设计了正负情感二分类实验,将选择的 8 000 条微博所标注的 7 种标签中的乐、好作为正向标签,哀、惧、怒、恶和惊作为负向标签,其中正、负向微博各 4 000 条。综合正负情感倾向值,得到的准确率见图 6。

通过图 6 得出,随着  $\lambda$  的增大,正负情感分类准确率提高,当  $\lambda = 0.4$  时准确率达到最高,当  $\lambda > 0.4$  时,正负情感分类准确率有下降的趋势,因此为了使颜文字在情感分类中更好地发挥作用,最终确定文本和 Emoji 组合、颜文字占比分别为 0.6 和 0.4。

#### 3.2.2 颜文字对最终情感分类准确率的影响

为了验证加入颜文字对情感分类准确度的影响,实验选取微博数据集 10 000 条,其中包含颜文

字的微博 8 000 条,不包含颜文字的微博 2 000 条,为了验证颜文字对微博情感分类准确率的影响,实验对比了 MNB 模型、ESM 模型、EMCNN 模型、DAM 模型以及本文提出的 EmotT 模型在 3 组数据集设置下的分类准确率:第 1 组全部微博共 10 000 条;第 2 组包含颜文字微博共 8 000 条;第 3 组在第 2 组数据集基础上去掉所有颜文字的微博 8 000 条。

从图 7、8 得出,含颜文字的情感分类性能指标始终高于不含颜文字的分类性能,表明分析颜文字的情感表达能力有助于提高情感分类准确率。此外,提出模型在 3 种数据集上的表现均优于其他模

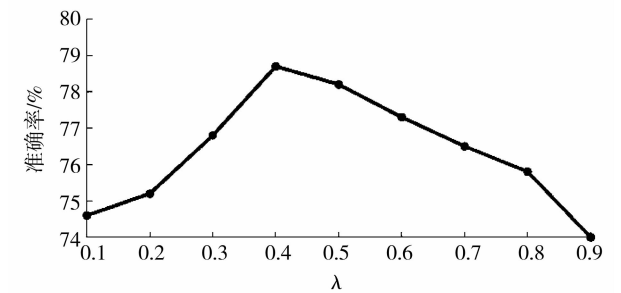


图6  $\lambda$  对情感判断准确率影响

Fig. 6 Influence of  $\lambda$  on the accuracy of sentiment classification

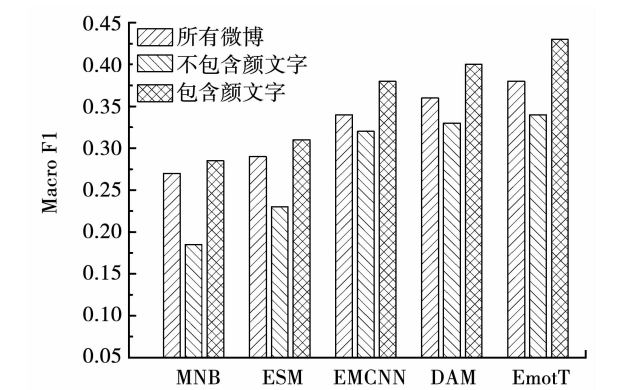


图7 情感分类结果 MacroF1

Fig. 7 Sentiment classification results of MacroF1

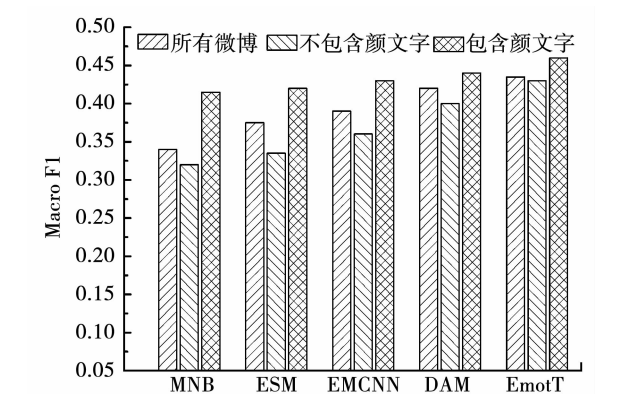


图8 情感分类结果 MicroF1

Fig. 8 Sentiment classification results of MicroF1

型,表明提出的针对短文本和 Emoji 组合部分情感分类模型可以更好地挖掘语句与情感标签之间的关联,加入的颜文字也进一步提高了情感分类性能。

## 4 结 论

本文提出了一种融合表情符号与短文本的多维情感分类方法,实现了对含有 Emoji 和颜文字的微博的情感分类,为群体情感细粒度分析提供了模型基础。主要贡献有:通过对日语颜文字的处理,构建中文颜文字词典;通过对 NLPCC2014 标准数据集的分析,本文模型得出的 MicroF1 和 MacroF1 值较其他模型分别平均提高 7.6% 和 6.6%,即该方法可以有效提高对文本和 Emoji 的多情感分类性能,验证了所提出的基于注意力机制的 CNN 和 LSTM 组合模型对微博文本和 Emoji 进行语义建模的有效性;对自行爬取的含有颜文字、Emoji 的微博语句进行多维情感分类,通过实验分析当文本、Emoji 组合部分与颜文字所占比例分别为 0.6 和 0.4 时所得到的多维情感分类准确率最高,且实验得出含颜文字的微博语句的 MicroF1 和 MacroF1 指标始终高于不含颜文字的微博语句,表明分析颜文字的情感表达能力有助于提高微博的情感分类准确率。通过以上结论,进一步说明本文提出的加入颜文字特征进行多维情感分类,提升了微博情感分类的性能,为中文社交媒体情感分析提供了新方法。

## 参考文献

- [1]王垒,李炳灿,唐楠棋,等. 95 后手机使用心理与行为白皮书[EB/OL]. (2019-04-18). [2019-7-4]  
WANG Lei, LI Bingcan, TANG Nanqi, et al. Psychological and behavioral report on mobile phone use after 1995[EB/OL]. (2019-04-18). [2019-7-4] <http://www.psy.pku.edu.cn/xwzx/xyxw/317324.htm>
- [2]新浪微博数据中心, 2018 微博用户发展报告[EB/OL]. (2019-03-15). [2019-7-4]  
Sina Weibo Data Center. 2018 Report on Weibo user development[EB/OL]. (2019-03-15). [2019-7-4]. <https://data.weibo.com/report/reportDetail? id=433>
- [3]READ J. Using emoticons to reduce dependency in machine learning techniques for sentiment classification[C]//Proceedings of the ACL Student Research Workshop. Ann Arbor, Michigan: ACL, 2005: 43. DOI: 10.3115/1628960.1628969
- [4]YANG Changhua, LIN K H Y, CHEN H H. Emotion classification using web blog corpora [C]// IEEE/WIC/ACM International Conference on Web Intelligence. Piscataway: IEEE, 2007: 275. DOI: 10.1109/WI.2007.51
- [5]SONG Kaisong, FENG Shi, GAO Wei, et al. Build emotion lexicon from microblogs by combining effects of seed words and emoticons in a heterogeneous graph [C]//Proceedings of the 26th ACM Conference on Hypertext & Social Media. Guzelyurt, Cyprus: ACM, 2015: 283. DOI: 10.1145/2700171.2791035
- [6]JIANG Fei, LIU Yiqun, LUAN Huanbo, et al. Microblog sentiment analysis with emoticon space model[J]. Journal of Computer Science and Technology, 2015, 30(5): 1120. DOI: 10.1007/s11390-015-1587-1
- [7]何炎祥,孙松涛,牛菲菲,等. 用于微博情感分析的一种情感语义增强的深度学习模型[J]. 计算机学报, 2017, 40(4): 773  
HE Yanxiang, SUN Songtao, NIU Feifei, et al. A deep learning model enhanced with emotion semantics for microblog sentiment analysis[J]. Chinese Journal of Computers, 2017, 40(4): 773. DOI: 10.11897/SP.J.1016.2017.00773
- [8]张仰森,郑佳,黄改娟,等. 基于双重注意力模型的微博情感分析方法[J]. 清华大学学报:自然科学版, 2018, 58(2): 122  
ZHANG Yangsen, ZHENG Jia, HUANG Gaijuan, et al. Microblog sentiment analysis method based on a double attention model[J]. Journal of Tsinghua University (Science & Technology), 2018, 58(2): 122. DOI: 10.16511/j.cnki.qhdxxb.2018.22.015
- [9]UTSU K, SAITO J, UCHIDA O. Sentiment polarity estimation of emoticons by polarity scoring of character components [C]//Proceedings of 2018 IEEE Region Ten Symposium. Sydney, NSW, Australia: IEEE, 2019: 237. DOI: 10.1109/TENCONSpring.2018.8691984
- [10]YU Chuanming, ZHU Xingyu, FENG Bolin, et al. Sentiment analysis of Japanese tourism online reviews[J]. Journal of Data and Information Science, 2019, 4(1): 89. DOI: 10.2478/jdis-2019-0005
- [11]YU Shuo, ZHU Hongyi, JIANG Shan, et al. Emoticon analysis for Chinese social media and e-commerce: The AZEemo system [J]. ACM Transactions on Management Information Systems, 2019, 9(4): 16. DOI: 10.1145/3309707
- [12]PTASZYNSKI M, MACIEJEWSKI J, DYBALA P, et al. CAO: A fully automatic emoticon analysis system based on theory of kinesics [J]. IEEE Transactions on Affective Computing, 2010, 1(1): 46. DOI: 10.1109/T-AFFC.2010.3
- [13]BIRDWHISTELL R L. Kinesics and context: Essays on body motion communication [M]. Philadelphia: University of Pennsylvania Press, 2010: 227
- [14]徐琳宏,林鸿飞,潘宇,等. 情感词汇本体的构造[J]. 情报学报, 2008, 27(2): 180  
XU Linhong, LIN Hongfei, PAN Yu, et al. Constructing the affective lexicon ontology [J]. Journal of the China society for Scientific and Technical Information, 2008, 27(2): 180. DOI: 10.3969/j.issn.1000-0135.2008.02.004
- [15]YAMADA T, TSUCHIYA S, KUROIWA S, et al. Classification of facemarks using N-gram [C]//Proceedings of International Conference on Natural Language Processing and Knowledge Engineering. Beijing, China: IEEE, 2007: 322. DOI: 10.1109/NLPKE.2007.4368050
- [16]BAHDANAU D, CHO K, BENGIO Y. Neural machine translation by jointly learning to align and translate [C]//Proceedings of 3rd International Conference on Learning Representations. San Diego, CA, United States: ICLR, 2015

(编辑 苗秀芝)