

基于情感语义对抗的跨语言情感分类模型^{*}

赵亚丽^{1,2}, 余正涛^{1,2}, 郭军军^{1,2}, 高盛祥^{1,2}, 相 艳^{1,2}

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

摘 要:传统的基于机器翻译的跨语言情感分类方法, 由于受机器翻译性能影响, 导致越南语等低资源语言的情感分类准确率较低。针对源语言和目标语言标记资源不平衡的问题, 提出一种基于情感语义对抗的跨语言情感分类模型。首先, 将句子和句子中情感词进行拼接, 用卷积神经网络对拼接后的句子分别进行特征抽取, 分别获得单语语义空间下的情感语义表征; 其次, 通过对抗网络, 在双语情感语义空间将带标签数据与无标签数据的情感语义表征进行对齐; 最后, 将句子与情感词最显著的表征进行拼接, 得到情感分类结果。基于汉英公共数据集和自主构建的汉越数据集的实验结果表明, 所提模型相比跨语言情感分类主流模型, 实现了双语情感语义对齐, 可以有效提升越南语情感分类的准确率, 且在差异性不同的语言对上也具有明显优势。

关键词:情感语义表征; 双语词嵌入; 低资源语言; 跨语言情感分类

中图分类号:TP391.41

文献标志码:A

doi:10.3969/j.issn.1007-130X.2023.02.017

A cross-language sentiment classification model based on emotional semantic confrontation

ZHAO Ya-li^{1,2}, YU Zheng-tao^{1,2}, GUO Jun-jun^{1,2}, GAO Sheng-xiang^{1,2}, XIANG Yan^{1,2}

(1.Faculty of Information Engineering and Automation, Kunming University of Science and Technology, Kunming 650500;

2.Artificial Intelligence Key Laboratory of Yunnan Province, Kunming 650500, China)

Abstract: Traditional cross-language sentiment classification methods based on machine translation are affected by the performance of machine translation, resulting in lower accuracy of sentiment classification in low-resource languages such as Vietnamese. Aiming at the problem of imbalance between source language and target language markup resources, this paper proposes a cross-language sentiment classification model based on sentiment semantic confrontation. Firstly, the sentences and the emotional words in the sentences are spliced, and the spliced sentences are jointly represented by the convolutional neural network, and the emotional semantic representations in the monolingual semantic space are obtained respectively. Secondly, through the confrontation network, the emotional semantic representations of labeled data and unlabeled data are aligned in the bilingual emotional semantic space. Finally, the most significant representations of sentences and emotional words are spliced together to obtain the results of emotional orientation classification. The experimental results based on the Chinese-English public data set and the Chinese-Vietnamese data set we constructed show that, compared with the mainstream methods of cross-language sentiment classification, the proposed method achieves bilingual sentimental semantic alignment, and can effectively improve the accuracy of sentimental orientation analysis of Vietnamese. The proposed method has obvious advantages in different language pairs.

^{*} 收稿日期:2021-04-26;修回日期:2021-09-05

基金项目:国家自然科学基金(61972186, 61762056, 61732005, 61761026);国家重点研发计划(2018YFC0830105, 2018YFC0830101, 2018YFC0830100);云南省高新技术产业专项(201606);云南省重大科技专项(202002AD080001-5);云南省基础研究计划(202001AT070047)

通信地址:650500 云南省昆明市昆明理工大学信息工程与自动化学院

Address: Faculty of Information Engineering and Automation, Kunming University of Science and Technology, Kunming 650500, Yunnan, P.R.China

Key words: emotional semantic representation; bilingual word embedding; low-resource language; cross-language sentiment classification

1 引言

随着“一带一路”倡议的提出,中国和越南在政治、经济等领域的联系愈加紧密。在越南热点事件发生后,越南社交媒体在互联网上进行报道,及时了解这些热点报道中民众的情感态度,对中越双边政治、贸易和文化交流有巨大影响。然而,越南语属于小语种,缺乏大规模标记数据,汉越 2 种语言存在语义鸿沟,且人工标注费时费力。

如表 1 所示,汉语和越南语属于不同的语言,对于同一领域的评论语句来说,汉语评论情感分类可以很容易判别,而用传统机器翻译方法翻译的越南语评论,可能存在极大的语义偏差。例如“Virus quá khủng khiếp để nói về nó”本意为:病毒令人谈之色变,而翻译结果为“病毒的变色”,导致无法直接判断越南语评论的情感分类。虽然已经获得 2 种语言的情感词典,但是中文有标记数据,而越南语完全没有标记数据,如何在越南语只有情感词典却没有任何标签的情况下,去实现越南语的情感分类,则成为了本文研究的核心难点。

Table 1 An example of the COVID-19 reviews in Vietnamese and Chinese

表 1 新冠疫情评论汉越样例数据示例

语言	评论句	情感词	有无标签
越南语	Virus rất khủng khiếp và sẽ mang đến cho chúng ta một cuộc khủng hoảng!	khủng khiếp	无
汉语	中国真棒,为全体医护人员点赞	棒 赞	有

越南语等低资源语言缺乏标注数据,而已有的监督学习方法仍然是基于机器翻译^[1-3]的跨语言情感分类方法。由于受机器翻译性能误差累积的影响,面向东南亚小语种语言的情感分类准确率普遍较低。

对于没有标注数据的低资源语言的情感分类,通常的做法是借助另外一种标记资源丰富的语言,用跨语言情感分类 CLSC (Cross-Lingual Sentiment Classification) 的方法辅助目标语言进行情感分类,借助对抗网络,获得双语语义对齐空间。传统的方法仅仅是在语义空间对齐,而情感词是情感分类的最直接有效的描述,也应该将其作为情感分类的依据。

2 相关工作

在跨语言情感分类的任务中,如何实现不同语言的句子在公共语义空间的语义对齐,如何在没有任何标签的情况下去无监督地完成情感分类是 CLSC 的核心难点。

要完成越南语的情感分类,首先要解决汉语和越南语不在同一语义空间的问题。Zhou 等^[4]实现跨语言情感分析的方式是通过机器翻译将源语言翻译为目标语言,跨语言表示学习是指不同语言的词向量表示可以共享一个向量空间,不同语言中情感语义相近的词在该空间中的距离相近。Mikolov 等^[5]提出将双语单词进行对齐,并训练得到源语言词向量空间到目标语言词向量空间的线性映射。Faruqui 等^[6]提出将源语言和目标语言的词嵌入映射到同一个向量空间。Lauzy 等^[7]提出通过自编码器对源语言进行编码,同时源语言和目标语言通过解码来得到双语的词向量。Meng 等^[8]利用平行语料库提升词典覆盖率,采用最大似然值对词语进行标注,进而提升情感分类的准确率。栗雨晴等^[9]通过构建双语词典,进行微博多情感分析。但这 2 种方法都需要构建多语言平行语料库,分类准确率依赖于语料库的质量和规模大小。Wang 等^[10]利用因子图模型的属性函数从每个帖子中学习单语和双语信息,利用因子函数来探索不同情绪之间的关系,并采用置信传播算法来学习和预测模型。

深度学习中的生成对抗网络可以很好地应用于迁移学习任务中,使用生成对抗网络构建起源语言与目标语言之间的桥梁,从而缓解了目标语言中标注数据匮乏的问题。但是,传统的 GAN (Generative Adversarial Networks)^[11]存在不易训练、生成数据可解释性差和易崩溃等缺点。近年来对 GAN 的研究主要有以下几种改进:对抗式自编码器 AAE (Adversarial AutoEncoders),由 Makhzani 等^[12]提出,加入自编码器可使生成数据更接近于输入数据,从而避免无效数据的产生;信息生成对抗网络 InfoGAN (interpretable representation learning by Information maximizing Generative Adversarial Net),Chen 等^[13]提出了 InfoGAN 模型,该模型在生成器中引入隐含编码,利用隐含编码对生成数据做出解释;序列生成对抗网络 Seq-

GAN, Yu 等^[14]提出的 SeqGAN 模型把序列生成问题看作序列决策制定过程,并使用强化学习的思想对模型进行训练;条件生成对抗网络 CGAN (Conditional GAN), CGAN 模型由 Mirza 等^[15]提出,该模型在对生成器和判别器建模时引入了条件变量,通过最大最小化条件变量使得生成器的输出既与真实数据相似又受条件约束。这些改进的 GAN 在文本和图像领域已取得不错的进展^[16]。

以上对于 GAN 的研究多数局限在图像以及单语文本研究领域。本文在基于对抗的卷积神经网络 CNN (Convolution Neural Network) 文本分类模型上加入了原文情感词典的特征扩展,经过对抗,将高资源情感语义模型迁移到低资源语言,得到双语情感语义对齐空间,显著提升低资源语言情感分类任务的识别性能。

3 基于情感语义对抗的跨语言情感分类模型

本文融合情感词典来指导低资源语言的跨语言情感分类,提出了基于情感语义对抗的无监督跨语言情感分类模型 SADAN (Sentiment Adversarial Deep Averaging Network)。图 1 是本文模型的基本结构,其中虚线表示无标签的数据,实线表示有标签的数据。

3.1 模型框架

SADAN 是一个具有 2 个分支的前馈网络。

网络中有 3 个主要模块,分别是:

(1) 融合情感词的句子语义表征模块:将句子和句子中的情感词进行拼接,对拼接后的句子进行嵌入,用 CNN 进行特征抽取,分别获得汉越 2 种语言在单语语义空间下的情感语义表征。CNN 作为共享特征提取器,目的是帮助情感分类器学习情感特征,并阻碍语言鉴别器辨别该特征来自于源语言还是目标语言。

(2) 双语情感语义对齐模块:将 2 种语言的语义表征进行对抗训练,实现高资源语言到低资源语言的对齐,得到双语情感语义对齐空间。

(3) 跨语言情感分类模块:将卷积得到的句子表征和句子中所包含的情感词(保持维度相同)进行拼接得到新的词向量,实现句子的特征扩展,再通过滤波器抽取得到句子特征的向量化表示。最后基于 softmax 激活函数^[18]进行目标语言的情感分类。

3.2 融合情感词的句子语义表征模块

3.2.1 情感词拼接

句子中情感词的获取是跨语言情感分类任务的第 1 步。利用匹配算法将语料句子中的每个词和情感词典中的词进行匹配,将句子中的情感词拼接在句子后面。

$$X = \{X_i, i = 1, 2, 3, \dots, n\} \quad (1)$$

$$X' = \{(X_i, S_i), i = 1, 2, 3, \dots, n; S_i \in \mathbf{R}^{n \times |S|}\} \quad (2)$$

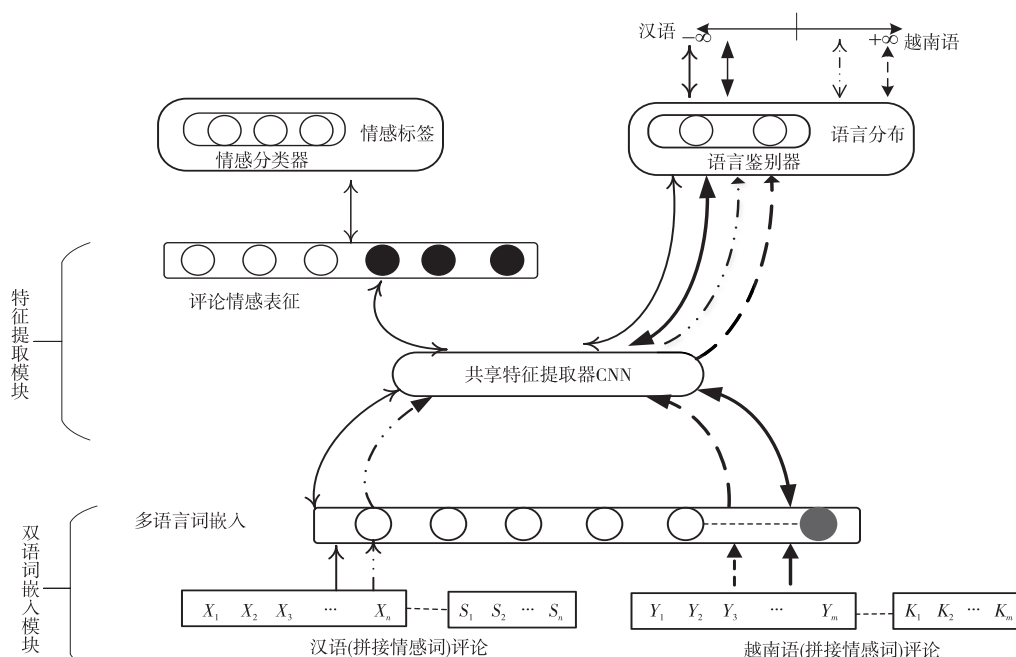


Figure 1 Cross-language sentiment classification model based on sentiment semantic confrontation

图 1 基于情感语义对抗的跨语言情感分类模型

$$Y = \{Y_j, j = 1, 2, 3, \dots, m\} \quad (3)$$

$$Y' = \{(Y_j, \mathbf{K}_j), j = 1, 2, 3, \dots, m; \mathbf{K}_j \in \mathbf{R}^{n \times |\mathbf{K}|}\} \quad (4)$$

其中, X 表示源语言句子集合, Y 表示目标语言句子集合, X' 表示源语言句子与句子中情感词拼接后的句子集, X_i 表示源语言句子集中的第 i 个句子, $\mathbf{S}_i$ 表示拼接在第 i 个句子后的情感词, n 表示源语言句子个数, $|\mathbf{S}|$ 表示拼接的情感词的长度, Y' 表示目标语言句子与句子中情感词拼接后的句子集, Y_j 表示目标语言句子中第 j 个句子, $\mathbf{K}_j$ 表示第 j 个目标语言句子后拼接的情感词, $|\mathbf{K}|$ 表示拼接的情感词的长度。

3.2.2 双语词嵌入

给定源语言句子输入 $X' = \{(x_i, l_i), j = 1, 2, 3, \dots, n\}$ 和一个目标语言句子输入 $Y' = \{(y_j, z_j), j = 1, 2, 3, \dots, m\}$ 。本文利用双语词嵌入将每个句子中的每个词表示成 z 维词向量, 如式(5)和式(6)所示:

$$\mathbf{E}_{X'} = \text{emb}(X_i, \mathbf{S}_i) \quad (5)$$

$$\mathbf{E}_{Y'} = \text{emb}(Y_j, \mathbf{K}_j) \quad (6)$$

其中, $\mathbf{E}_{X'} \in \mathbf{R}^{n \times |q|}$ 和 $\mathbf{E}_{Y'} \in \mathbf{R}^{m \times |d|}$ 分别表示嵌入函数, 它将每一个输入序列中的每个词转化为对应的 z 维词向量; $|q|$ 和 $|d|$ 表示源语言和目标语言输入模型的句子长度。本文所采用的词嵌入设为 50 维, 即 $z = 50$ 。

3.3 双语情感语义对齐模块

将句子和句子中的情感词进行拼接, 用卷积神经网络对拼接后的句子进行联合表征, 分别获得汉越 2 种语言单语语义空间下的情感语义表征; 然后, 通过对抗网络在双语情感语义空间将带标签数据与无标签数据的情感语义表征进行对齐。该模块包括共享特征提取器模块(F)和语言鉴别器(D)。

3.3.1 CNN 共享特征提取器(F)

图 1 中嵌入层的输出送给用于特征提取的卷积层。本文在对句子进行特征提取的同时, 也对句子中的情感词进行特征提取。每个卷积层都有固定大小的滑动窗口, 每次只处理窗口内的信息。窗口的大小设定为 k , 在卷积操作中有连续 k 个词向量获得新的特征向量 $\mathbf{c}_i$, i 表示第 i 个特征值, $\mathbf{x}_{i:i+k-1}$ 表示输入评论句中第 i 个词到第 $i+k-1$ 个词经过卷积操作得到的向量表示。操作过程可以用式(7)表示:

$$\mathbf{c}_i = f_1(\mathbf{w} \cdot \mathbf{x}_{i:i+k-1} + \mathbf{b}) \quad (7)$$

其中, 滤波器的权重矩阵 $\mathbf{w} \in \mathbf{R}^{k \times d}$, $\mathbf{b}$ 为偏置项,

f_1 为激活函数。

提取出来的特征 $\mathbf{C}$ 表示为式(8):

$$\mathbf{C} = [\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_{n-k+1}] \quad (8)$$

同理, 情感词抽取出的特征 $\mathbf{A}$ 表示为式(9):

$$\mathbf{A} = [\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_m] \quad (9)$$

其中, $\mathbf{a}_i, 1 \leq i \leq m$ 表示情感词抽取出的特征向量。

3.3.2 语言鉴别器(D)

考虑源语言(src)和目标语言(tgt)的联合隐藏特征的分布, 如式(10)和式(11)所示。

$$P_F^{\text{src}} \triangleq P(F(x) | x \in S) \quad (10)$$

$$P_F^{\text{tgt}} \triangleq P(F(x) | x \in T) \quad (11)$$

为了学习汉越双语的语言特征, SADAN 训练 $F(x)$ 使这 2 个分布尽可能接近, 以获得更好的跨语言概括。由于 Jensen-Shannon 散度存在不连续性, Arjovsky 等^[19] 建议将 Wasserstein 距离最小化, 并证明了其对超参数选择的稳定性有所帮助。

在本文的 SADAN 模型中, 根据 Cédric 等^[20] 提出的 Kantorovich Rubenstein 对偶性原理, 最小化了 P_F^{src} 和 P_F^{tgt} 之间的 Wasserstein 距离 W 。操作过程可以用式(12)表示:

$$W(P_F^{\text{src}}, P_F^{\text{tgt}}) = \sup_{g: \mathbb{R}^d \rightarrow \mathbb{R}} E_{f(x) \sim P_F^{\text{src}}} [g(f(x))] - E_{f(x') \sim P_F^{\text{tgt}}} [g(f(x'))] \quad (12)$$

其中, 上确界(最大值)取所有 1-Lipschitz 函数 g 的集合。为了近似地计算 $W(P_F^{\text{src}}, P_F^{\text{tgt}})$, 本文使用语言鉴别器 D 作为式(12)中的函数 g , 目标是寻求式(12)中的上确界。 D 试图为源语言实例输出较高的分数, 为目标语言输出较低的分数。

为了使 D 成为 Lipschitz 函数(直到一个常数), D 的参数总是被限制在一个固定的范围内。设 D 用 θ_q 参数化, 那么目标 J_q 如式(13)所示:

$$J_q(\theta_f) \equiv \max_{\theta_q} E_{F(x) \sim P_F^{\text{src}}} [D(F(x))] - E_{F(x') \sim P_F^{\text{tgt}}} [D(F(x'))] \quad (13)$$

其中, J_q 是 P_F^{src} 和 P_F^{tgt} 之间 Wasserstein 距离的近似值。

3.4 跨语言情感分类模块(P)

最大值池化操作的实际作用就是保留某个滤波器提取到的最显著的特征。最大池化层可以减少训练模型的参数规模, 还可以过滤掉一些不必要的噪声。通过多个不同大小的滤波器生成的特征

值 $\hat{C}$ 进行组合获得分类特征 $\mathbf{v}$ 。经过全连接层的操作可以将特征进一步量化,从而抽取到更深层的语义特征 $\mathbf{v}'$,量化的过程用式(14)表示如下:

$$\mathbf{v}' = \mathbf{w}' \cdot \mathbf{v} + \mathbf{b}' \quad (14)$$

其中, $\mathbf{w}'$ 为全连接层训练的权重矩阵, $\mathbf{b}'$ 为偏置项。 $\mathbf{v}'$ 经过全连接层,获得了多个特征类别范围内的估计值,需要做归一化的处理,采用 softmax 分类函数可以决策出最大概率的类别,用式(15)表示如下:

$$p = \text{softmax}(\mathbf{v}') \quad (15)$$

其中, softmax 为分类器。 p 表示句子最终情感特征所属的概率,可以判别出句子的情感类别。

对于由 θ_p 参数化的情感分类器 P ,使用传统的交叉熵损失,表示为 $L_p(y', y)$,其中 y' 和 y 分别是预测的标签分布和真实的标签。 L_p 是 P 预测正确标签的负对数似然。因此,求 P 的以下损失函数的最小值,用式(16)表示如下:

$$J_q \equiv \min_{\theta_q} E [L_p(P(F(x)), y)] \quad (16)$$

4 实验和结果分析

4.1 数据集和评价指标

4.1.1 汉越数据集

汉越数据集的构建过程与文献[13]中构建 CLSC 汉英数据集类似,参数设置如表 2 所示。CLSC 数据集是从新浪微博与推特平台获取的 2020 年美国疫情相关的热门社交媒体评论数据,共包含 20 334 条中文微博评论和 11 233 条越南语推特评论,经过筛选与预处理形成 json 格式文件,经过一系列数据整理和预处理后获得汉越 CLSC 数据集,数据格式为:(中文句子,中文情感词,标签 l)和(越南语句子,越南语情感词),其中 $l \in \{0, 1, 2, 3, 4\}$ 。

4.1.2 英中数据集

本文的英语数据集是文献[21]中的 70 万条 Yelp 评论的平衡数据集并采用了他们的训练集和验证集分割:65 万份用于训练和 50 万份用于验证。中文数据集方面,本文使用 Lin 等[22]的 1 万条中国酒店评论用作验证集,另外的 15 万条未标记的中国酒店评论作为测试集。

4.2 评价指标及模型参数设置

4.2.1 评价指标

本文的实验评价指标使用准确率 $Accuracy$ 、精确率 $Precision$ 、召回率 $Recall$ 和 $F1$ 值,主要使用 $Accuracy$ 进行评价。其计算公式分别如式

(17)~式(20)表示为:

$$Accuracy = \frac{(TP + TN)}{\text{总样本}} \quad (17)$$

$$Precision = \frac{TP}{TP + FP} \times 100\% \quad (18)$$

$$Recall = \frac{TP}{TP + FN} \times 100\% \quad (19)$$

$$F1 = \frac{2 \times P \times R}{P + R} \times 100\% \quad (20)$$

其中, TP 表示实际为正例且被分类器划分为正例的样本数, TN 表示实际为负例且被分类器划分为负例的样本数, FP 表示实际为正例且被分类器划分为负例的样本数, FN 表示实际为负例且被分类器划分为正例的样本数。

4.2.2 模型参数设置

本文模型基于 pytorch 的深度学习框架,具体参数设置如表 2 所示。

Table 2 Parameters setting

表 2 参数设置

模型参数	参数值
优化器	Adam
模型维度	50
Max-epoch	30
batch_size	100
dropout	0.2
learning rate	5e-5
Kernel_num	400
Sequence_Kernel_size	[3,4,5]
Word_Kernel_size	[1,1,1]
Kernel_num	400

4.3 基准模型与实验结果分析

4.3.1 基准模型

本文模型与 Train-on-SOURCE-only, Domain Adaptation, Machine Translation, CLD-based CLTC 和 ADAN 等基准模型做了对比实验。

Train-on-SOURCE-only^[23]模型: Logistic Regression 和 DAN 在源语言英语上进行训练,并且只依靠双语词嵌入 BWE(Bilingual Word Embeddings)对目标语言进行分类。

Domain Adaptation 模型: 在域自适应中,广泛使用 Sinno 等^[24]的模型并不奏效,因为它需要样本数量(650 000)的二次空间。因此,本文将其与 Chen 等^[25]提出的 mSDA 模型相比较,后者是对亚马逊评论进行跨领域情感分类非常有效的模型。

基于 Machine Translation CLSC^[26] 模型:根据机器翻译的 2 种模型 Logistic Regression+MT 和 DAN+MT 评估本文模型。

CLD-based CLTC(Cross-Lingual Distillation-based Cross-Lingual Text Classification)模型: Xu 等^[27] 提出了一种跨语言提取 CLD 模型,该模型利用并行语料库上的预测来训练目标语言 CLD-KC-NN(Cross-Lingual Distillation-Knowledge-aware Convolutional Neural Network),并进一步提出了一种改进的变体 CLDFA-KCNN(Cross-Lingual Distillation with Feature Adaptation-Knowledge-aware Convolutional Neural Network),该变体利用对抗性训练来弥补源语言和目标语言中标记和未标记文本之间的领域差距。

ADAN (Adversarial Deep Averaging Network)模型^[28]:该模型由 CNN 和 GAN 神经网络组成,其中 CNN 负责提取句子中的特征,GAN 负责学习双语语言特征。

4.3.2 实验结果及分析

实验对本文所提模型的有效性进行验证,在越南语数据集上证明模型的有效性。实验结果如表 3 所示。

Table 3 Experimental results on Chinese-Vietnamese data set

表 3 汉越数据集上的实验结果 %	
评价指标	结果
Accuracy	40.23
Precision	40.06
Recall	42.77
F1	40.34

为了验证本文所提出的基于情感语义对抗的无监督跨语言情感分类模型的泛化性,本文还在文献[13]中公布的公共数据集上将性能最优的“本文模型”与上述基准模型的性能作对比,对比实验结果如表 4 所示。实验表明了情感语义对抗的有效性与泛化性。具体分析结果如下:

(1) Train-on-SOURCE-only 基准模型中的 Logistic Regression 使用标准的监督学习算法,此外,本文还评估了模型的一个非对抗变量,即 DAN 模型,它是情感分类的现代神经模型之一。与 SADAN 相比,仅基于源语言的基线模型表现不佳,这表明 BWE 本身不足以转移知识。

(2)mSDA 的表现并不具有竞争力,这可能是因为在包括 mSDA 在内的许多领域适应模型都是为使用词袋特征而设计的,但这种模型并不适合本文

的任务,因为 2 种语言的词汇完全不同。这表明即使是域自适应算法也不能在 CLSC 任务中使用现成的 BWE。

(3)SADAN 模型在 2 种语言上都显著优于机器翻译基准模型,这表明本文的对抗模型可以在没有任何目标语言注释数据的情况下成功地进行跨语言情感分类。

(4)与 CLD-based CLTC 基准模型相比较,可以看出本文 SADAN 模型使中文准确率有了显著提升,CLD-based CLTC 模型使用对抗式训练在单一语言中进行领域适应,而本文直接使用对抗式训练进行跨语言概括比较,证明了本文 SADAN 模型的有效性。

(5)ADAN 模型仅仅得到语义对齐,而本文模型得到双语情感语义对齐,证明了本文模型的有效性。

由此可见,本文的分类模型相对其他基准模型具有更高的准确率。

Table 4 Accuracy comparison of models on yelp dataset and Chinese hotel dataset

表 4 yelp 和中文酒店数据集上模型准确率对比 %

模型		准确率
Train-on-SOURCE-only	Logistic Regression	30.58
	DAN	29.11
Domain Adaptation	mSDA	31.44
MT_CLSC	Logistic Regression+MT	34.01
	DAN + MT	39.66
CLD-based CLTC	CLD-KCNN	40.96
	CLDFA-KCNN	41.82
ADAN		42.49
本文模型		45.65

4.3.3 拼接是否为情感词对模型准确率的影响

为了证明本文提出的融合情感词典来指导跨语言情感分类对本文模型的有效性,本文在汉越和中英数据集上进行了一组简单的消融实验,实验结果如表 5 所示,“(−)情感词”表示未使用情感词辅助指导跨语言情感分类,仅使用语句中的情感无关随机词。

Table 5 Ablation experimental results about the emotional words for splicing

表 5 拼接是否为情感词消融实验结果 %

模型	准确率	
	汉越数据集	中英数据集
(−)情感词	36.74	40.07
(本文模型)情感词	40.23	45.65

由表 5 可知,在本文模型中使用情感词拼接时,较无情感词拼接的模型准确率提高了 4%。由此证明,在跨语言任务当中,拼接情感词可以丰富短文本的表征,从而提高了模型的准确率。

4.3.4 有无对抗网络对模型准确率的影响

为了证明本文提出的使用对抗网络辅助来指导跨语言情感分类方法对本文模型的有效性,本文在数据集上进行了是否应该有对抗网络的消融实验,实验结果如表 6 所示,“(−)对抗”表示未使用对抗网络辅助来指导跨语言情感分类。

Table 6 Ablation experimental results about the against network

表 6 是否有对抗网络消融实验结果 %		
模型	准确率	
	汉越数据集	中英数据集
(−)对抗	37.09	39.98
本文模型对抗	40.23	45.65

由表 6 可知,在本文模型中,当使用对抗网络时,准确率较高。由此证明,在跨语言任务当中,对抗也可以在一定程度上拉近不同语言的语义空间。

4.3.5 拼接情感词长度选择对模型准确率的影响

为了证明拼接不同长度情感词来指导跨语言情感分类对本文模型的有效性,本文在汉越和中英数据集上进行了拼接情感词长度选择的消融实验,实验结果如表 7 所示,“(mean)情感词平均个数”表示使用训练中情感词的平均长度进行补齐来指导跨语言情感分类,“(max)情感词个数”表示使用训练中情感词的最大长度进行补齐来指导跨语言情感分类。

Table 7 Results of the ablation experiment on the number of emotional words

表 7 情感词个数消融实验结果 %		
模型	汉越数据集 准确率	中英数据集 准确率
(mean)情感词平均个数	40.23	45.65
(max)情感词最大个数	39.77	43.78

由表 7 可知,在本文模型中,在原来评论句子的基础上,加入情感词后,准确率较基准 ADAN 模型有一定提升,说明情感词拼接的个数可以影响模型的准确率,当情感词拼接个数取该批次中情感词长度的平均值时,模型效果达到最佳。

情感词拼接个数的增加会使情感相关特征得到扩展,从而达到更好的分类效果。而情感词增加到该批次的最大数量时,准确率开始逐渐下降。英

文 Yelp 的评论数据和中文酒店相关的评论数据截然不同,但在拼接了该批次的情感词平均长度后,模型准确率也有明显提升,准确率达到 45.65%;当拼接的情感词数量增加到批次长度限制的时候,准确率有显著下降。这说明随着拼接情感词个数增加,卷积层从情感词和评论句子拼接后的向量中学习到的特征会更分散,这时候情感词的加入反倒产生了噪声,导致准确率在后续不再增长。因此,拼接的情感词个数不是越多越好。

5 结束语

本文研究旨在提升无标注的低资源目标语言的情感分类的准确率。针对不同语言之间存在语义鸿沟等导致分类准确率低这一问题,本文利用有丰富标记数据的源语言辅助无标注数据的目标语言,提出基于情感语义对抗的跨语言情感分类模型。本文模型在自制的汉越数据集和 CLSC 公共数据集上均取得了显著效果。当前,针对低资源语言的跨语言情感分析仍然是情感分析领域的研究热点和难点。在未来工作中,将针对低资源语言信息检索展开进一步研究,在提高分类准确率的同时,在细粒度方面展开进一步研究。

参考文献:

- [1] Hermann K M, Blunsom P. Multilingual models for compositional distributed semantics[C]// Proc of the 52nd Annual Meeting of the Association for Computational Linguistics, 2014:58-68.
- [2] Gouws S, Bengio Y, Corrado G, Bilbowa. Fast bilingual distributed representations without word alignments[C]// Proc of the 32nd International Conference on Machine Learning, 2015:748-756.
- [3] Zhou H, Chen L, Shi F, et al. Learning bilingual sentiment word embeddings for cross-language sentiment classification [C]// Proc of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing, 2015:430-440.
- [4] Zhou X, Wan X, Xiao J. Cross-lingual sentiment classification with bilingual document representation learning[C]// Proc of the 54th Annual Meeting of the Association for Computational Linguistics, 2016:1403-1412.
- [5] Mikolov T, Le Q V, Sutskever I. Exploiting similarities among languages for machine translation [J]. arXiv: 1309. 4168, 2013.
- [6] Faruqui M, Dyer C. Improving vector space word representations using multilingual correlation[C]// Proc of the European Chapter of the Association for Computational Linguistics, 2014:462-471.

- [7] Lauly S, Larochelle H, Khapra M. An autoencoder approach to learning bilingual word representations[J]. Advances in Neural Information Processing Systems, 2014, 3: 1853-1861.
- [8] Meng X, Wei F, Liu X, et al. Cross-lingual mixture model for sentiment classification[C]//Proc of the Meeting of the Association for Computational Linguistics, 2012: 572-581.
- [9] Li Yu-qing, Li Xin, Han Xu, et al. A bilingual lexicon-based multi-class semantic orientation analysis for microblogs[J]. Acta Electronica Sinica, 2016, 44(9): 2068-2073. (in Chinese)
- [10] Wang Z, Lee S Y M, Li S, et al. Emotion analysis in code-switching text with joint factor graph model[J]. IEEE/ACM Transactions on Audio Speech & Language Processing, 2017, 25(3): 469-480.
- [11] Wang Kun-feng, Gou Chao, Duan Yan-jie, et al. Generative adversarial networks: The state of the art and beyond[J]. Acta Automatica Sinica, 2017, 43(3): 321-332. (in Chinese)
- [12] Makhzani A, Shlens J, Jaitly N, et al. Adversarial autoencoders[J]. arXiv:1511.05644, 2015.
- [13] Chen X, Duay Y, Houthoof R, et al. Infogan: Interpretable representation learning by information maximizing generative adversarial nets[C]//Proc of International Conference on Neural Information Processing Systems, 2016: 2172-2180.
- [14] Yu L, Zhang W, Wang J, et al. SeqGAN: Sequence generative adversarial nets with policy gradient[C]//Proce of the 31st AAAI Conference on Artificial Intelligence, 2017: 2852-2858.
- [15] Mirza M, Osindero S. Conditional generative adversarial nets[J]. arXiv:1411.1784, 2014.
- [16] Spurra A, Aksane E, Hilliges O. Guiding InfoGAN with semi-supervision[C]//Proc of Joint European Conference on Machine Learning and Knowledge Discovery in Databases, 2017: 119-134.
- [17] Chen Wen-bing, Guan Zheng-xiong, Chen Yun-jie. Data augmentation method based on conditional generative adversarial net model[J]. Journal of Computer Applications, 2018, 38(11): 3305-3311. (in Chinese)
- [18] Tang Xian-lun, Du Yi-ming, Liu Yu-wei, et al. Image recognition with conditional deep convolutional generative adversarial networks[J]. Acta Automatica Sinica, 2018, 44(5): 855-864. (in Chinese)
- [19] Arjovsky M, Chintala S, Bottou L. Wasserstein generative adversarial networks[C]//Proc of the 34th International Conference on Machine Learning, 2017: 214-223.
- [20] Cédric V, Ludger R. Optimal transport: Old and new[J]. Jahresbericht der Deutschen Mathematiker-Vereinigung, 2009, 111(2): 18-21.
- [21] Zhang X, Zhao J B, LeCun Y. Character-level convolutional networks for text classification[C]//Proc of Neural Information Processing Systems, 2015: 649-657.
- [22] Lin Y O, Lei H, Wu J, et al. An empirical study on sentiment classification of chinese review using word embedding[C]//Proc of the 29th Pacific Asia Conference on Language, Information and Computation, 2015: 258-266.
- [23] Kingma D, Ba J. Adam: A method for stochastic optimization[J]. arXiv:1412.6980, 2014.
- [24] Sinno J P, Ivor W T, James T, et al. Domain adaptation via transfer component analysis[J]. IEEE Transactions on Neural Networks, 2011, 22(2): 199-210.
- [25] Chen X, Sun Y, Athiwaratkun B, et al. Adversarial deep averaging networks for cross-lingual sentiment classification[J]. arXiv:1606.01614v4, 2016.
- [26] Carmen B, Rada M, Janyce W, et al. Multilingual subjectivity analysis using machine translation[C]//Proc of the Conference on Empirical Methods in Natural Language Processing, 2008: 127-135.
- [27] Xu R C, Yang Y M. Cross-lingual distillation for text classification[C]//Proc of the 55th Annual Meeting of the Association for Computational Linguistics, 2017: 1415-1425.
- [28] Chen X, Sun Y, Athiwaratkun B, et al. Adversarial deep averaging networks for cross-lingual sentiment classification[J]. Transactions of the Association for Computational Linguistics, 2018, 6: 557-570.

附中文参考文献:

- [9] 栗雨晴, 礼欣, 韩煦, 等. 基于双语词典的微博多类情感分析方法[J]. 电子学报, 2016, 44(9): 2068-2073.
- [11] 王坤峰, 苟超, 段艳杰, 等. 生成式对抗网络 GAN 的研究进展与展望[J]. 自动化学报, 2017, 43(3): 321-332.
- [17] 陈文兵, 管正雄, 陈允杰. 基于条件生成式对抗网络的数据增强方法[J]. 计算机应用, 2018, 38(11): 3305-3311.
- [18] 唐贤伦, 杜一铭, 刘雨微, 等. 基于条件深度卷积生成对抗网络的图像识别方法[J]. 自动化学报, 2018, 44(5): 855-864.

作者简介:



赵亚丽(1995-), 女, 山西临汾人, 硕士生, CCF 会员(G5333G), 研究方向为自然语言处理、深度学习和情感分析。E-mail: 610650183@qq.com

ZHAO Ya-li, born in 1995, MS candidate, CCF member(G5333G), her research interests include natural language processing, deep learning, and emotion analysis.



余正涛(1970-), 男, 云南曲靖人, 博士, 教授, CCF 会员(E20-0011683S), 研究方向为自然语言处理和信息检索与信息抽取。E-mail: ztyu@hotmail.com

YU Zheng-tao, born in 1970, PhD, professor, CCF member(E20-0011683S), his research interests include natural language processing, and information retrieval & extraction.