

北京大学学报(自然科学版)
Acta Scientiarum Naturalium Universitatis Pekinensis
doi: 10.13209/j.0479-8023.2016.004

基于词语情感隶属度特征的情感极性分类

宋佳颖 黄旭 付国宏[†]

黑龙江大学计算机科学技术学院, 哈尔滨 150080; [†] 通信作者, E-mail: ghfu@hotmail.com

摘要 在模糊集合论框架下探索基于词语情感隶属度的情感极性分类特征表示方法。首先以 TF-IDF 为权重分别构建情感特征词语的正向、负向极性隶属度, 并以隶属度对数比作为分类特征值构建基于支持向量机的情感极性分类系统。在产品评论、NLPCC2014 情感分类评测数据和 IMDB 英文影评等数据上的实验结果表明, 基于情感隶属度特征的系统优于基于布尔、频度和词向量等特征表示的系统, 从而验证了所提出的基于情感隶属度特征表示的有效性。

关键词 情感极性分类; 模糊集合论; 隶属度; 支持向量机

中图分类号 TP391

Exploiting Lexical Sentiment Membership-Based Features to Polarity Classification

SONG Jiaying, HUANG Xu, FU Guohong[†]

School of Computer Science and Technology, Heilongjiang University, Harbin 150080;
[†] Corresponding author, E-mail: ghfu@hotmail.com

Abstract A lexical sentiment membership based feature representation was presented for Chinese polarity classification under the framework of fuzzy set theory. To this end, the TF-IDF weighted words are first used to construct the corresponding positive and negative polarity membership for each feature word. Then, the log-ratio of each membership is computed. Finally, a support vector machines based polarity classifier is built with the membership log-ratios as its features. Furthermore, the classifier is evaluated over different datasets, including a corpus of reviews on automobile products, the NLPCC2014 data for sentiment classification evaluation and the IMDB film comments. The experimental results show that the proposed sentiment membership feature representation outperforms the state of the art feature representations such as the Boolean features, the frequent-based features and the word embeddings based features.

Key words sentiment polarity classification; fuzzy sets; membership; supported vector machines

随着 Web 2.0 的兴起和社会媒体的迅速发展, 情感分析(亦称意见挖掘)已成为自然语言处理研究的一个热点, 并在近年得到快速发展, 各种情感分析系统层出不穷。由于机器学习方法性能的不断提升, 使得情感分类能够得到效果较好的基线系统, 机器学习框架能够从选取的特征中学习不同类别的指向信息, 其参数、特征集和权重的确定对分类性能起决定性作用。因此, 本文将特征的选择和表示

作为重点探索的内容。随着基于神经网络的语言模型的发展, 基于词袋(bag-of-words)的模型逐渐受到排斥, 神经网络模型不再只是对词语的罗列而更多地考察了词序关系, 在大规模的无监督训练下往往能得到更多的语义信息, 因此在抽取、语音识别、翻译、校对等工作中发挥了优势。同时, 很多情感分析工作已将词语、段落的向量表示作为特征权重应用于分类框架^[1-4], 但通过对比发现, 这类方法在

情感分类问题的解决上依然有很大的改进空间。虽然 n-grams 语言模型考虑了词语的窗口内上下文, 由于很少有数据能够满足多窗口的上下文的短语概率计算, n-grams 存在数据稀疏和高维度的限制, 对于词语间的语义距离的衡量依然模糊。同时与 n-grams 相对的递归神经网络(recurrent neural networks, RNNs)语言模型^[2], 其内部结构能够更好地进行平滑预测, 从而放宽了上下文的窗口限制, 在很多应用中优于传统的 n-grams, 因此我们利用 RNNs 作为本文的一组基线方法。然而, 在新方法不断涌现的今天, 词袋模型仍然存在优势, 某些语料数据在传统的朴素贝叶斯(Naïve Bayes, NB)、支持向量机(Support Vector Machine, SVM)分类框架下, 结合优化的特征、权重集, 依然能够获得更好的分类结果^[5]。

本文从优化特征及权重的角度出发, 在已经取得很好效果的 NB-SVM 基础上, 进一步探索更佳的 SVM 应用方法。我们分别对产品评论语料应用递归神经网络语言模型^[2](Recurrent neural network based language model, RNNLM)通过贝叶斯法则判定正负极性, 应用 Paragraph Embedding 生成的句子向量作为特征通过 SVM 分类器判定极性^[3], 应用 NB-SVM^[5]结合 n-grams 特征判定极性作为本文的基线方法。受到情感表达外沿模糊性的启发, 我们尝试用模糊集合理论挖掘词语的正负情感间的细微差别, 结合模糊推理的词汇模糊集合的情感隶属度确定方法, 将正负情感极性隶属度有效融合作为特征表示方法, 提出基于词语情感隶属度特征的分类框架, 并完成与上述各种分类方法的比对, 实验结果说明了本文提出的分类方法对于情感极性分类的有效性。

1 相关研究

情感分析问题通常分为两大解决线路, 分别是基于词典的情感分类方法^[6-7]和基于语料库的情感分类方法^[5,8-9]。由于通用词典对于各类型、领域的文本覆盖度不足, 基于词典的方法的效用逐渐弱化, 而基于对语料库中信息进行统计的机器学习方法越发受到重视。Yang 等^[9]将句子级情感分类看做序列标注问题, 有情感标签的句子作为输入, 通过条件随机场和后序正规化(posterior regularization)来学习参数, 利用上下文短句的语境和评价对象对不含有情感词语的短句进行情感倾向预测, 对各类特

征进行整合, 包括词典模板、转折连接、意见共指等。随着情感分析研究的不断深入以及基于语言模型的新型语义表示方法^[2,10]的出现, 各种基于神经网络模型的向量表示方法^[10-13]也被应用到情感分析领域。由于这些基于神经网络的语言模型能够在无监督的条件下挖掘一定的语义信息, 这些向量表示的获取也成为当前的研究热点。Le 等^[3]通过词语的向量表达预测上下文的词向量, 将句子向量看做是一个特殊的主题词向量, 应用随机梯度下降训练词语语义向量表示, 用这些词向量进一步推断句子向量, 将得到的向量表示作为支持向量机分类器的特征完成句子情感分类。Bespalov 等^[13]通过浅层语义分析得到词的向量表示, 进一步将文本表示为 n-grams 特征向量对应的线性权重向量用于情感分析。Tang 等^[14]在大规模微博语料库中以微博表情符号作为弱情感标签, 通过三种神经网络模型有监督的训练得到面向情感语义的词向量表示, 将词向量表示作为特征放入 SVM 分类器中, 得到不错的效果。Wang 等^[5]分别对朴素贝叶斯和 SVM 这两种常用的分类模型的适用情况进行分析, 提出应用朴素贝叶斯对数频次比作为 SVM 特征权重的分类模型, 通过实验验证了这种简单的模型对于情感分类任务十分有效。本文在 Wang 等^[5]工作的基础上, 以相关理论为依据进一步对特征权重优化, 以得到更佳分类性能。

2 情感分类方法

2.1 情感词语模糊集合

正向词语模糊集合: 设论域 X 为所有词语的集合, 则论域 X 上的正向情感词语模糊集合 POS 是 X 到 $[0,1]$ 的一个映射 $\mu_{\text{POS}}: X \rightarrow [0,1]$ 。对于 $x \in X$, μ_{POS} 称为正向词语模糊集合 POS 的隶属度函数, $\mu_{\text{POS}}(x)$ 称为 x 属于词语模糊集合 POS 的隶属度。

负向词语模糊集合: 设论域 X 为所有词语的集合, 则论域 X 上的负向情感词语模糊集合 NEG 是 X 到 $[0,1]$ 的一个映射 $\mu_{\text{NEG}}: X \rightarrow [0,1]$ 。对于 $x \in X$, μ_{NEG} 称为负向词语模糊集合 NEG 的隶属度函数, $\mu_{\text{NEG}}(x)$ 称为 x 属于词语模糊集合 NEG 的隶属度。

由上述定义可知, 这些隶属度函数的确定是模糊集合理论能否有效投入应用的关键。

2.2 情感分类的 TF-IDF

2.2.1 情感词语频率 TF

定义 $f^{(i)} \in R^{|V|}$ 是训练样例 i 的特征的频数, 即

有 $f^{(i)}$ 代表了特征 V_j 在样例 i 中出现的次数。对于所有的训练样例, 可以定义正负两类特征频数向量如式(1)和(2)。其中 α 是为了数据平滑设置的参数。

$$TF_{POS} = \alpha + \sum_{i: y^{(i)}=1} f^{(i)}, \quad (1)$$

$$TF_{NEG} = \alpha + \sum_{i: y^{(i)}=-1} f^{(i)}. \quad (2)$$

根据上面得到特征频数向量, 对 TF_{POS} 和 TF_{NEG} 分别除以其自身向量的频数总和进行归一化处理, 进一步计算其对数比, 如式(3)所示。

$$r = \log \left(\frac{TF_{POS} / \|TF_{POS}\|_1}{TF_{NEG} / \|TF_{NEG}\|_1} \right). \quad (3)$$

2.2.2 情感词语的逆文档频率 IDF

NB-SVM 是将文档词频信息的归一化对数比作为特征权重, 其形式如式(3)。受到基于模糊推理的词语隶属度构建方法的启发^[15], 我们通过分析认为, 在归一化频数的基础上, 融合特征对应各个类别的逆文档频率(IDF)信息, 能够使特征具有更好的类别指向性, 从而削弱在各类极性的情感句中都有大量出现的无关特征对分类性能的影响, 可以作为词语的模糊情感极性隶属度的一种表示方法。因此, 为词语计算对应的正负两类的 IDF_{POS} 和 IDF_{NEG} , 如式(4)和(5)所示。

$$IDF_{POS}^i = \log \left(\frac{s_{pos} + s_{neg}}{\text{Count}_{pos}^i} \right), \quad (4)$$

$$IDF_{NEG}^i = \log \left(\frac{s_{pos} + s_{neg}}{\text{Count}_{neg}^i} \right), \quad (5)$$

其中, Count_{pos}^i 表示含有特征 i 且极性为正向的样例的数量, 反之为负向, 计算时同样使用加 1 平滑。 s_{pos} 和 s_{neg} 分别表示训练数据中正向极性样例和负向极性样例的数量。

2.3 词语情感隶属度

常见的隶属度函数确定方法包括模糊统计法、例证法、专家经验法等。为了避免在选择时受到主观因素的过多影响, 本文采用模糊统计法计算每个词语的正、负情感隶属度。模糊统计法是通过 n 次重复独立统计实验来确定某个特征词对正、负情感词语模糊集合的隶属度, 其形式上与概率统计法比较类似, 但二者分别属于不同的数学模型。

我们以 TF-IDF 表示法为原型, 通过对频数向量的归一化能够平衡词频对极性类别的影响, 归一化向量对应的与相同极性的 IDF 的积做为每个特征对于正负情感极性的最终隶属度, 正负情感隶属度计算如式(6)和(7)。

$$M_{POS} = (TF_{POS} / \|TF_{POS}\|_1) IDF_{POS}, \quad (6)$$

$$M_{NEG} = (TF_{NEG} / \|TF_{NEG}\|_1) IDF_{NEG}. \quad (7)$$

2.4 词语情感隶属度特征表示

2.3 节定义了基于 TF-IDF 的词语情感隶属度函数, 能够给每个特征确定它隶属于两个情感极性模糊集合的程度。为了量化正负情感隶属度大小对特征的情感指向的作用, 我们将两类隶属度函数值进行融合, 把正负情感隶属度的对数比作为特征权重值, 特征 i 的权重计算方法如式(8)所示。

$$r_i = \log \left(\frac{(TF_{POS}^i / \|TF_{POS}\|_1) IDF_{POS}^i}{(TF_{NEG}^i / \|TF_{NEG}\|_1) IDF_{NEG}^i} \right). \quad (8)$$

2.5 支持向量机 SVM

对于支持向量机, 它的基本原理是通过对有类标记的训练数据构造相应的模型, 继而应用模型通过测试数据中的属性特征来预测其对应的类标记。训练数据形式是成对的样例和标签 (x_i, y_i) , $i=1, \dots, r$, 其中 $x_i \in \mathbb{R}^n$, $y_i \in \{-1, +1\}$, 为了解决某些样本点线性不可分, 引进了松弛变量 $\xi_i \geq 0$, 改变约束条件为 $y_i(w \cdot x_i + b) \geq 1 - \xi_i$, 目标函数由原来的 $\frac{1}{2} \|w\|^2$ 变成

$$\min_w \frac{1}{2} w^T w + C \sum_{i=1}^l \xi_i(w; x_i, y_i), \quad (9)$$

其中, $C > 0$ 是惩罚系数, 它决定了对于误分类的惩罚的大小, 一般根据实际问题确定。由于 Linear 是应对大规模训练任务的快捷有效的 SVM 分类器, 且 Linear 能够支持 L2-regularized 逻辑回归(LR)和 L2-loss、L1-loss 线性支持向量机, 选择 Linear^①作为本文的 SVM 工具, 可选训练参数 s 为 0, 即应用 L2 正规化逻辑回归, 对应的上式中 $\xi = \log(1 + e^{-y_i w^T x_i})$ 。

3 实验结果与分析

为了对上述方法进行全面的验证, 分别对汽车领域产品评论、NLPCC 评测^②的数据和英文影评

① <http://www.csie.ntu.edu.tw/~cjlin/liblinear>

② http://tcci.ccf.org.cn/conference/2014/pages/page04_eva.html

IMDB^①数据进行情感极性分类。本节给出相应的实验设置、结果及其分析。

3.1 实验设置

如表 1 所示, 我们给出三类实验数据的统计信息, 语料分别是从小汽车之家^②爬取的汽车领域的多品牌网络用户评价、NLPCC2014 评测中的情感分类任务数据(多领域产品评论)和 IMDB(大规模英文公开影评)。其中 IMDB 数据共有影评 10 万句, 使用方法与 Le 等^[3]相同, 包含有标注的 25000 条训练语句、25000 条测试语句, 其余 5 万句是无标注的语句, 仅在无监督地训练词向量时使用, 标注的语句分为正向极性、负向极性两类标签。实验的评测指标为准确率(accuracy, Acc)、精确率(Precision, P)、召回率(Recall, R)和 F -测度(F)。

为了进一步验证基于情感隶属度的特征表示的有效性, 本文还考虑用以下 4 种方法作为实验的基线方法。

1) RNNLM + NaïveBayes: Mikolov 等^[2]提出的基于递归神经网络的语言模型(RNNLM)在语音识别实验的结果中验证了 RNNLM 明显优于 n -gram 语言模型。此处 RNNLM 基于简单的 Elman 神经网络^[16], 它是一个包含输入层、隐藏状态层和输出层的神经网络, 能够允许应用更大的窗口的上下文来完成对序列中其他词的预测, 在训练时做到了更好的数据平滑。但在实际训练中, 上下文的窗口的大小还会受梯度下降效率的限制。本文利用 RNN 语言模型, 借助贝叶斯法则计算每个测试样例属于正负极性类别的概率, 从而完成分类。本文 RNNLM 相关实验应用 RNNLM Toolkit^③完成, 具体训练参数设定为 `-hidden(50)`, `-direct-order(3)`, `-direct(200)`, `-class(100)`, `-debug(2)`, `-bptt(4)`, `-bptt-block(10)`。

表 1 语料统计信息

Table 1 Basic statistics of the experiment data

项目	汽车评论		NLPCC2014		IMDB	
	正向	负向	正向	负向	正向	负向
训练	1103	1103	5000	5000	12500	12500
测试	1194	1110	1250	1250	12500	12500

① <http://ai.stanford.edu/~amaas/data/sentiment/>

② <http://www.autohome.com.cn/>

③ <http://www.fit.vutbr.cz/~imikolov/rnnlm/>

④ <https://code.google.com/p/word2vec/>

2) Paragraph Vector + SVM: Le 等^[3]提出的无监督的对句子、段落或文本预测得到定长的向量表示可以作为特征用于有监督的分类框架。具体的, 将句子向量看做是一个特殊的主题词向量, 应用随机梯度下降训练词语语义向量表示, 再用这些词向量进一步推断句子向量表示, 将得到的向量表示作为支持向量机分类器的特征完成句子情感分类。其中, 句子向量合成的相关实验借助 word2vec^④完成。在训练句子向量阶段, 我们选择的语言模型为 Skip-Gram, 向量维度设定了不同的大小(100, 200 和 300), 训练的窗口大小设定为 10, 同时使用 NEG 和 HS 方法, 其他参数为默认的值。

3) Bool + SVM: 最为传统的布尔权重支持向量机应用同样作为本文的基线系统实验, 分别考察不同特征集结合布尔权重的分类效果。

4) NB-SVM: 由 Wang 等^[5]提出的线性分类器, 它是由归一化特征频数的对数比作为特征权重的基于支持向量机的分类框架。

为了全面对比特征与特征权重的结合对分类效果的影响, 选择在相关研究中常用的有效的类别指向信息^[3,5]作为本文的特征集: 1) 基于 n -grams 的特征集, 包含一元语法词组(uni-gram)、二元语法词组(bigram)和三元语法词组(trigram); 2) 基于词性信息的特征集, 包括名词、动词、形容词、代词、数词、量词等实词。由于否定副词和一些程度副词也是对情感表达有指向作用的词汇, 本文实词特征还加入副词特征。

3.2 实验结果与分析

3.2.1 汽车评论语料情感极性分类结果

针对汽车产品评论, 设置实验及其结果如表 2 所示, 在 Paragraph Vector 相关实验中, 鉴于对生成的语义向量表示的准确性的考虑, 在无监督的向量训练阶段, 我们在训练语料中加入 26729 句爬取得到的网络汽车评论作为背景语料, 帮助得到更为有效的 embedding 向量表示。在生成句子向量表示时, 分别考察了不同维度大小对结果的影响, 表 2 第一列括号里的数字表示生成的向量的维数。本文提出的将词语情感隶属度对数比作为特征权重的方法, 在实验结果中以 fuzzy + SVM 作为标记。

从表 2 可以看出,在特征选择方面,通常三元语法特征会优于二元语法特征,二元语法特征会优于一元语法特征,但在 SVM 结合布尔权重和应用 NB-SVM 时却不符合我们的理论推断。分析原因为语料规模较小,数据稀疏带来的结果的不稳定性;另外,简单的布尔权重使得大部分三元特征的权重为 1,无法很好地衡量这些多词组特征的情感指向比重。而在句子向量(Paragraph vector)和情感隶属度对数比特征的 SVM (Fuzzy + SVM)实验结果中,特征不同时呈现的分类性能都符合常规的理论推断,在一定程度说明了三元语法特征较二元、一元特征具有更好的限定性,能够更准确地获取句中的词序关系。同时,从准确率方面来看,虽然实词特征较一元的词语特征更为有效,但依然不如二元、三元短语特征,说明高阶的短语特征使组合的词语具有更准确的限定性,更全面涵盖句子情感信息。在分类效果方面,可以看出原有方法中的 NB-SVM 具有较好的分类性能,随着特征的优化能够得到更佳的结果,同时其结果优于基于 RNN 语言模型和句子向量合成的方法,说明虽然语义向量信息的获取能够促进抽取、相似度衡量等工作的发展,但如何从语义信息中有针对性的挖掘情感信息,仍有待研究。本文提出的 fuzzy+SVM 在同等特征集作用时,取得了优于 NB-SVM 的分类效果,进一步说明

在确定特征权重时,在特征频数归一化的基础上,融合 IDF 信息后,去除了在正负极性中都大量出现的特征对隶属度的影响,使得到的特征情感隶属度更能全面描述各个特征对于类别的指向作用。

3.2.2 NLPCC2014 评测数据情感极性分类结果

为了进一步验证方法的性能,首先使用 NLPCC 评测的公开数据进行实验,本轮实验主要考察性能较好且比较接近的三类基于支持向量机的方法,由于数据规模的限制会很大程度的影响无监督训练的过程,本轮实验没有采用训练句子向量作为特征,表 3 列出同样使用 NLPCC 数据的 Wang 等^[17]的结果用于比对。

从表 3 可以看出: 1) 同类方法不同特征的对比下,呈现出三元语法特征优于二元语法特征,而二元语法特征也好于一元语法特征现象,这完全符合高阶语法模型能够更准确地限定上下文的特点,同时反映出语料规模较小时(汽车评论),对理论的验证可能存在偏差,容易对研究方法的走向形成错误指引; 2) 在 NLPCC 数据集上的实验结果表明,基于情感隶属度对数比特征的系统在所有评测指标中均取得最好性能。表 3 中, Wang 等^[17]采用的是通过深度学习得到的词语向量特征表示结合逻辑回归分类器的方法。NLPCC2014 评测数据集上的对比实验结果表明本文提出的基于隶属度的特征表示方

表 2 汽车评论情感极性分类结果
Table 2 Results of polarity classification over car reviews

方法	Acc	P_{pos}	R_{pos}	F_{pos}	P_{neg}	R_{neg}	F_{neg}
RNNLM + Naïve Bayes	73.52	77.55	68.84	72.94	70.10	78.56	74.09
Bool + SVM (uni-gram)	79.30	81.70	77.39	79.48	76.98	81.35	79.11
Bool + SVM (bigram)	79.30	82.38	76.38	79.27	76.44	82.43	79.32
Bool + SVM (trigram)	78.99	82.16	75.96	78.94	76.08	82.25	79.05
Paragraph Vector+SVM(100)	79.21	83.32	74.87	78.87	75.63	83.87	79.54
Paragraph Vector+SVM (200)	79.34	83.12	75.46	79.10	75.98	83.51	79.57
Paragraph Vector+SVM (300)	79.81	82.99	76.80	79.77	76.90	83.06	79.86
NB-SVM (unigram)	80.77	84.11	77.55	80.70	77.72	84.23	80.85
NB-SVM (bigram)	81.51	85.56	77.38	81.27	77.94	85.95	81.75
NB-SVM (trigram)	81.29	86.23	76.05	80.82	77.14	86.94	81.75
Fuzzy + SVM (实词)	81.25	83.13	80.07	81.57	79.38	82.52	80.92
Fuzzy + SVM (unigram)	81.12	83.97	78.56	81.18	78.43	83.87	81.06
Fuzzy + SVM (bigram)	81.46	85.35	77.55	81.26	78.01	85.68	81.67
Fuzzy + SVM (trigram)	81.68	86.62	76.47	81.23	77.52	87.30	82.12

说明: 粗体数字表示该指标下的最好结果。以下同。

表 3 NLPCC2014 评测数据集上的情感分类结果
Table 3 Result of polarity classification over NLPCC2014

方法	Acc	P_{pos}	R_{pos}	F_{pos}	P_{neg}	R_{neg}	F_{neg}
Bool + SVM (uni-gram)	75.44	75.00	76.32	75.65	75.90	74.56	75.22
Bool + SVM (bigram)	77.16	76.75	77.92	77.33	77.58	76.40	76.99
Bool + SVM (trigram)	77.72	77.39	78.32	77.85	78.06	77.12	77.59
NB-SVM (bigram)	78.08	77.86	78.48	78.17	78.31	77.68	77.99
NB-SVM (trigram)	78.92	78.76	79.20	78.98	79.08	78.64	78.86
Wang 等 ^[17]	—	78.00	74.80	76.40	75.80	78.90	77.30
Fuzzy + SVM (实词)	76.32	75.99	76.96	76.47	76.66	75.68	76.17
Fuzzy + SVM (bigram)	78.32	78.05	78.80	78.42	78.59	77.84	78.22
Fuzzy + SVM (trigram)	79.36	79.22	79.60	79.41	79.50	79.12	79.31

表 4 IMDB 数据集上的情感分类结果
Table 4 Result of polarity classification over the IMDB dataset

方法	Acc	P_{pos}	R_{pos}	F_{pos}	P_{neg}	R_{neg}	F_{neg}
RNNLM + Naïve Bayes	86.28	86.29	86.28	86.29	86.27	86.28	86.27
Paragraph Vector+SVM (300)	89.31	89.30	89.31	89.31	89.31	89.30	89.31
Bool + SVM (trigram)	90.17	89.60	90.89	90.24	90.76	89.46	90.10
NB-SVM (trigram)	91.87	90.97	92.98	91.97	92.82	90.77	91.78
Fuzzy + SVM (trigram)	91.93	90.96	93.10	92.02	92.94	90.75	91.83

法的有效性。

3.2.3 IMDB 情感极性分类结果

除中文产品评论和 NLPCC2014 评测数据集以外,我们还选择了常用于情感分类任务的英文语料 IMDB 数据,并且应用各类方法的最好参数进行情感分类,包含代表性最强的 trigram 特征以及语义表示效果最好的 300 维向量特征。在完成句子向量特征的实验 Paragraph Vector 时,我们在无监督训练阶段没有借助其他数据,而是使用完整的 IMDB 数据共 100000 句训练得到对应的句子向量。实验结果如表 4 所示。

从表 4 可以看出,在 IMDB 数据集上的实验结果中,本文方法得到的综合准确率和 F 值都表现出最大优势,精确率和召回率均处于较好位置,说明本文确定的情感隶属度是对词语极性和强度的有效度量。在 Wang 等^[5]的工作中针对 IMDB 数据得到 91.22% 的准确率,相比之下,本文提出的基于词语情感隶属度的特征值表示方法更具有实际意义。由

于本文方法完全是基于语料库的统计方法,不对语言种类、领域有任何限定,上述结果中的英文数据实验就形成了本文方法有效性的完整印证。

4 结论与展望

根据情感极性分类研究现状,本文在现有方法的基础上,以 TF-IDF 为原型,融合模糊推理的隶属度确定方法,进一步为词语设定了情感极性隶属度,从而得到基于词语情感隶属度的特征值表示方法。分别对汽车领域评论、NLPCC 评测数据和 IMDB 数据集进行实验,结果显示通过优化特征和权重,在传统的机器学习分类框架下依然能够取得很好的分类性能。

虽然本文实验取得了预期结果,证明了融合的情感隶属度特征值对于情感分类问题的有效性,但没能在整体框架下实现全面创新,仅取得小幅度的提高。后续工作应该全面深化对问题的研究,扩大数据规模并挖掘更有效的有指向性的特征。

参考文献

- [1] Socher R, Pennington J, Huang E H, et al. Semi-supervised recursive autoencoders for predicting sentiment distributions // Proceedings of EMNLP'11. East Stroudsburg, 2011: 151–161
- [2] Mikolov T, Karafiát M, Burget L, et al. Recurrent neural network based language model // Proceedings of INTERSPEECH'10. Chiba, 2010: 1045–1048
- [3] Le Q V, Mikolov T. Distributed representations of sentences and documents. arXiv preprint arXiv:1405.4053, 2014: 1188–1196
- [4] Zhang Dongwen, Xu Hua, Su Zengcai, et al. Chinese comments sentiment classification based on word2vec and SVM perf. Expert Systems with Applications, 2015, 42(4): 1857–1863
- [5] Wang S, Manning C D. Baselines and bigrams: simple, good sentiment and topic classification // Proceedings of ACL'12. Jeju Island, 2012: 90–94
- [6] Ding Xiaowen, Liu Bing, Yu P S. A holistic lexicon-based approach to opinion mining // Proceedings of WSDM'08. New York, 2008: 231–240
- [7] Taboada M, Brooke J, Tofiloski M, et al. Lexicon-based methods for sentiment analysis. Computational Linguistics, 2011, 37(2): 267–307
- [8] Wang Hongning, Lu Yue, Zhai Chengxiang. Latent aspect rating analysis on review text data: a rating regression approach // Proceedings of SIGKDD'10. New York, 2010: 783–792
- [9] Yang Bishan, Claire Cardie. Context-aware learning for sentence-level sentiment analysis with posterior regularization // Proceedings of ACL'14. Baltimore, 2014: 325–335
- [10] Mikolov T, Chen K, Corrado G, et al. Efficient estimation of word representations in vector space. arXiv preprint arXiv:1301.3781, 2013
- [11] Socher R, Perelygin A, Wu J Y, et al. Recursive deep models for semantic compositionality over a sentiment treebank // Proceedings of EMNLP'13. Seattle, 2013: 1631–1642
- [12] Bengio Y, Courville A, Vincent P. Representation learning: a review and new perspectives. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2013, 35(8): 1798–1828
- [13] Bessalov D, Bai B, Qi Y, et al. Sentiment classification based on supervised latent n-gram analysis // Proceedings of CIKM'11. Glasgow, 2011: 375–382
- [14] Tang Duyu, Wei Furu, Yang Nan, et al. Learning sentiment-specific word embedding for twitter sentiment classification // Proceedings of ACL'14. Baltimore, 2014: 1555–1565
- [15] Aida-zade K, Rustamov S, Mustafayev E, et al. Human-computer dialogue understanding hybrid system // Proceedings of the 2012 International Symposium on Innovations in Intelligent Systems and Applications (INISTA). Trabzon, 2012: 1–5
- [16] Elman J L. Distributed representations, simple recurrent networks, and grammatical structure. Machine Learning, 1991, 7(2/3): 195–225
- [17] Wang Yuan, Li Zhaohui, Liu Jie, et al. Word vector modeling for sentiment analysis of product reviews // CCIS(NLPCC'14). Shenzhen, 2014, 496: 168–180