

DOI: 10.13718/j.cnki.xdzk.2022.11.022

基于混合优化的双模深度学习文本分类方法

吴绪玲

西南交通大学 希望学院 信息工程系, 成都 610400

摘要: 为解决文本分类中分类精度低的问题, 提出一种混合优化的双模深度学习文本分类方法。该方法设计了一种混合优化算法, 对深度学习模型进行权值调优, 得到相关度高的特征和高性能文本分类结果。首先对文档进行预处理得到特征集合, 设计了基于乌鸦搜索算法(CSA)和蝗虫优化算法(GOA)的混合优化算法, 并使用双向门控循环单元(GRU)进行特征选择, 得到具有上下文语义信息且相关的特征。最后, 将最优特征输入到混合优化的深度置信网络(DBN)中得到文本分类结果。

关键词: 文本分类; 混合优化算法; 深度学习; 双向门控循环

单元; 深度置信网络

中图分类号: TP391

文献标志码: A

开放科学(资源服务)标识码(OSID):



文章编号: 1673-9868(2022)11-0234-09

Dual-Mode Deep Learning Text Classification Method Based on Hybrid Optimization

WU Xuling

School of Information Engineering, Hope College, Southwest Jiaotong University, Chengdu 610400, China

Abstract: In order to solve the problem of low classification accuracy in text classification, a dual-mode deep learning text classification method based on hybrid optimization is proposed. This method designs a hybrid optimization algorithm, optimizes the weights of the deep learning model, and obtains optimal features and high-performance text classification results. First, the document is preprocessed to obtain the feature set, and a hybrid optimization algorithm based on the crow search algorithm and the grasshopper optimization algorithm is designed. The hybrid optimized Bi-GRU is used to select the optimal features, and the highly relevant features with context semantic information are obtained. Finally, the optimal features are input into the DBN with hybrid optimized weights to obtain the text classification results.

收稿日期: 2021-09-23

基金项目: 四川省电子商务与现代物流研究中心项目(DSWL21-27); 成都市哲学社会科学重点研究基地课题(2019002)。

作者简介: 吴绪玲, 硕士, 讲师, 主要从事软件工程方面的研究。

Key words: text classification; crow search algorithm; grasshopper optimization algorithm; bi-gated recurrent unit; deep confidence network

作为信息检索、信息过滤等领域的基础，文本分类技术起着至关重要的作用，成为自然语言处理的研究热点^[1-2]。文本分类的效率和准确度影响了文本处理能力，现在对海量文本的有效处理为语义分析等打下良好基础^[3]。文本分类中分类结果取决于特征质量及分类方法。如今，电子文档的数量越来越多，文档特征也随之增加，然而，大多数提取的特征部分是不相关且多余的，高维特征会增加文本分类的计算时间，降低分类结果的准确性^[4]。

特征选择是文本分类预处理中克服这一问题的重要步骤，其目的是确定一个包含有限数量的相关特征的子集，这些特征足以维持或提高分类任务的性能^[5]。传统的特征选择方法有信息增益、卡方分布和主成分法等。文献[6]提出一种文本分类的特征选择方法，设计了类间集中度和类内分散度评估函数进行特征选择，文献[7]提出了改进卡方度量的特征选择方法，但在大数据时代，这些传统方法很难保证文本分类的准确性。文献[8]提出了使用优化算法的特征选择，使用蚁群算法和人工神经网络算法实现了用于文本分类的特征选择过程，但是特征选择没有考虑到上下文信息和相关性，且收敛速度慢，易发生停滞，连续优化问题求解较弱。

近些年，将选择的特征子集进行分类得到最终的分类结果，由机器学习发展到深度学习，都已经被用于文本分类。常用的文本分类方法有决策树、支持向量机、K 最近邻居、朴素贝叶斯和神经网络等^[9]。文献[10]提出一种 CNN-ELM 混合短文本分类模型，该方法卷积神经网络提取特征并进行优化，然后使用误差最小化极速学习机作为分类器完成短文本分类任务；文献[11]提出了多项式和伯努利朴素贝叶斯的文本分类方法，虽然分类生成的解更简单有效，但是某个参数的估计和新数据的到达会带来很多麻烦，如需要计算先验概率、分类决策存在错误率、对输入数据的表达形式很敏感等；文献[12]提出了注意力机制和卷积层的双向长短期记忆网络(Long Short-Term Memory, LSTM)的文本分类方法，包含双向 LSTM，关注机制和卷积层，得到上下文相关的特征，最后 softmax 分类器用于对处理后的上下文信息进行分类；文献[13]提出一种卷积递归神经网络的文本分类方法，所提出的分类模型使用卷积神经网络来提取本地特征，然后利用 LSTM 网络的巨大内存能力连接提取的特征，提高文本分类精度；文献[14]中引入胶囊网络替代 CNN 提取文本特征，将其与 LSTM 连接形成融合神经网络模型对文本进行分类。文献[10-14]中基于深度学习的特征学习方法中，对于深度模型的参数寻优方面的研究较少。

为了加大深度学习文本分类准确率，提出一种混合优化的双模深度学习文本分类方法，设计混合优化算法，对双模深度学习模型权值寻优，使用双向 GRU 进行特征选择，将选择的上下文相关特征子集输入优化后的深度信念网络，从而实现文本分类，提高分类准确率等性能。

1 深度学习模型

1.1 双向 GRU

循环神经网络(Recurrent Neural Network, RNN)具有短时记忆，RNN 通过反向传播算法将误差往前传递。当输入序列较长时，RNN 存在梯度爆炸或消失的问题，LSTM 和 GRU(图 1)的引入来解决循环神经网络这些问题。LSTM 作为主流 RNN 模型常应用于文本信息提取，优点是有效保留文档上下文时序关系；而 GRU 将 LSTM 模

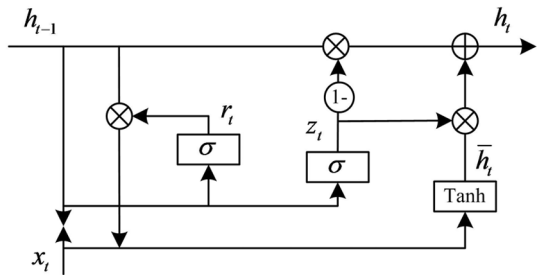


图 1 GRU 结构

型进行简化,只有更新门和重置门两种,以更少的参数和计算量进行信息特征提取^[15].

更新门有助于确定需要将多少过去的信息传递给模型,使用此功能,可以防止模型复制过去的所有信息;重置门用于决定有多少过去的信息在模型中被遗忘和确定.更新门和重置门如式(1)所示^[2].

$$\begin{cases} z_t = \sigma(W_z h_{t-1} + x_t U_z) \\ r_t = \sigma(W_r h_{t-1} + x_t U_r) \end{cases} \quad (1)$$

其中, σ 表示激活函数, x_t 表示输入向量, h_{t-1} 表示前一时间的输出, W_z, U_z, W_r 和 U_r 表示权值.当前隐藏层内容 $\bar{h}_t$ 经过重置门处理,与当前的输入相结合,并通过 Tanh 激活函数得到.最终的 GRU 输出结果为^[2]:

$$h_t = z_t \times \bar{h}_t + (1 - z_t) h_{t-1} \quad (2)$$

在双向 GRU 中,模型以向前和向后的顺序处理输入序列,将来自两个方向的知识合并到输出中.

1.2 深度置信网络

与传统的机器学习模型相比,DBN(图 2)是一个概率生成模型,通过对神经网元权重的训练,使得神经网络以最大概率生成训练数据.DBN 是由多层受限玻尔兹曼机(Restricted Boltzmann Machine, RBMs)堆叠而成,RBM 由基于加权连接的隐藏单元和可见单元组成,多层网络具有解决任何复杂任务的能力,从而使数据分类更有效地确定增量式文本分类^[16].

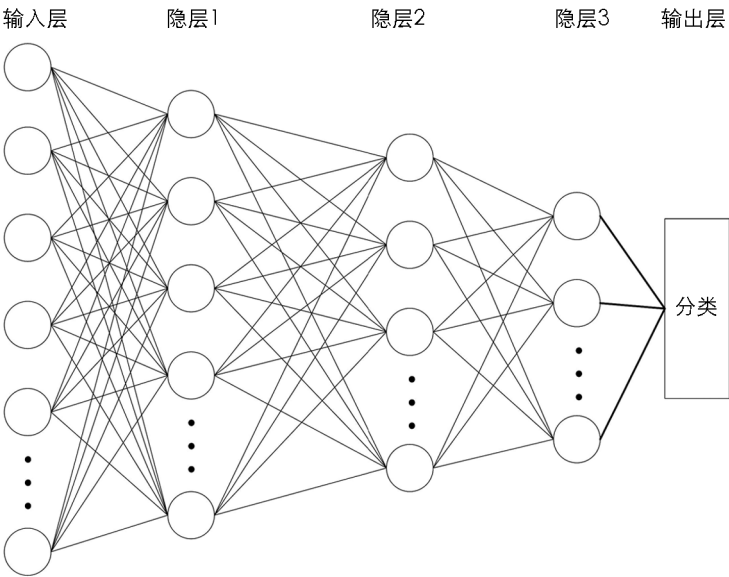


图 2 DBN 结构

DBN 是由 RBMs 堆叠得到的,隐藏层的能量函数表示为:

$$E(v, h) = - \sum_{i=1}^p \sum_{j=1}^q v_j w_{ij} h_i - \sum_{j=1}^q b_j v_j - \sum_{i=1}^p a_i h_i \quad (3)$$

其中, v_j 表示第 j 个可见节点, h_i 表示第 i 个隐藏节点, w_{ij} 表示两个节点间的权值, a_i 和 b_j 分别表示隐藏层和可见层的偏置.

在 DBN 学习中,随机进行参数初始化,将导致局部最优问题,所以需要进一步地学习来优化网络参数.本文设计了乌鸦搜索算法(Crow Search Algorithm, CSA)和蝗虫优化算法(Grasshopper Optimization Algorithm, GOA)的混合优化算法,实现对 DBN 的权值优化,解决局部最优问题,使得特征选择和分类性能提高.

2 基于混合优化算法的双模深度学习文本分类

将文档中的关键字输入预处理阶段,以便使用停止词删除和词干处理过程,从数据中消除不一致和冗

余的词,之后将文本进行特征表示,然后设计混合优化算法,对深度学习模型进行权值优化,使用双向 GRU 进行特征选择,最后使用优化 DBN 实现有效的文本分类.

2.1 文本文档的预处理

为了将文本转换为数字形式,使用了预处理措施.首先进行分词,将一个句子或者序列分解成单个的单词和标点.然后去停止词,去除文档中如虚词、介词和连词等特征词,减少特征总数;之后进行词干提取,通过消除每个单词的前缀和后缀,词干分离为特定单词的相关形式,赋予了几乎相似的词根,在本文中使用波特词干分析法进行分析.

将文本形式的数据转换为数字形式是词加权过程,为了测量文档表示的词权重,使用了词频-逆文档频率(Term Frequency-Inverse Document Frequency, TFIDF),每个文档的权重向量如式(4)所示^[5].

$$\begin{cases} T_i = (tw_{u1}, tw_{u2}, \dots, tw_{uv}, \dots, tw_{ur}) \\ tw_{uv} = xf(u, v) * yf(u, v) \end{cases} \quad (4)$$

其中,文本文件 u 中的词 v 权重表示 tw_{uv} ,文本文件 u 中词 v 出现的频率表示 $xf(u, v)$.通过式(4)对词进行特征统计,实现文本中词的特征表示.

2.2 混合优化的双模深度学习模型

通过将乌鸦搜索算法(Crow Search Algorithm, CSA)引入蝗虫优化算法(Grasshopper Optimization Algorithm, GOA),生成混合优化算法,对双模深度学习模型进行训练,实现特征选择和文本分类.混合优化算法继承了两种优化算法的优点,并提供更好的增量文本分类性能.

CSA 算法是基于乌鸦的智能行为而设计的一种元启发式算法,该算法用于控制算法的多样性,更容易实现,具有较低的收敛性和对超参数高度敏感的特点. GOA 完全基于蝗虫群的行为,猎物的搜索是通过一系列步骤(例如包围、开发和探索)在搜索空间内或搜索空间外进行的, GOA 能够解决未知搜索空间的实际问题.在 GOA 中引入 CSA,跳出局部最优,提供更好的收敛速度,同时获得全局最佳解决方案.

以随机方式初始化深度学习模型的权重,并表示如下:

$$Y = \{Y_1, Y_2, \dots, Y_g, \dots, Y_a\} \quad (5)$$

其中 α 表示总的权重.每当将新文档添加到网络时,都会计算误差并更新权重,如果计算的误差小于先前文档的评估误差,则根据混合优化算法更新权重,更新后的权重是通过计算存储的权重与范围度之间的差来计算的.

$$Y(e+1) = Y(e) \pm R$$

其中 $Y(e)$ 表示存储的权重, R 是范围,可表示为

$$R = \sqrt{d(\log 2s/d + 1) - \log(\delta/4)/s}$$

其中 d 表示特征集的 Vapnik-Chervonenkis 维度,用于描述学习机实现的一组特征的能力, s 表示训练样本数, δ 表示 $[0, 1]$ 之间的随机数.

根据 GOA 优化,得到:

$$Y_h^i = c \left(\sum_{k \neq h}^{k=1} c \frac{B_{ui} - B_{li}}{2} f(p_k^i - p_h^i) \frac{p_k - p_h}{D_{h,k}} \right) + \hat{U}_i \quad (6)$$

其中, c 表示递减系数, B_{ui} 表示第 i 维的上限, B_{li} 表示第 i 维的下限, f 表示受其他蝗虫的影响力函数, p_k^i 表示第 i 个维度中第 k 个蝗虫的位置, p_h^i 表示第 i 个维度中第 h 个蝗虫的位置, $D_{h,k}$ 表示第 h 和 k 个蝗虫的距离, $\hat{U}_i$ 表示前进的趋势.基于 CSA 的权重更新是基于搜索食物的概率.

$$Y_h(j+1) = Y_h(j) + r_h \times L_h(j) \times (m_k(j) - Y_h(j)) \quad (7)$$

其中 $Y_h(e)$ 表示乌鸦在当前迭代 j 中的位置, r_h 表示随机数, $L_h(j)$ 表示当前迭代中第 h 维的飞行长度, $m_k(j)$ 表示乌鸦的记忆.将式(8)进行简化得到:

$$m_k(j)=\frac{1}{r_h\times L_h(j)}[Y_h(j+1)+Y_h(j)(r_h\times L_h(j)-1)]$$

(8)

将 CSA 引入到 GOA, 即 $\hat{U}_i=m_k(j)$, 带入公式(6)得到:

$$Y_h^i=c\left(\sum_{k\neq h}^k c\frac{B_{ui}-B_{li}}{2}f(p_k^i-p_h^i)\frac{p_k-p_h}{D_{h,k}}\right)+$$
$$\frac{[Y_h(j+1)+Y_h(j)(r_h\times L_h(j)-1)]}{r_h\times L_h(j)}$$

(9)

则混合优化算法的优化问题求解如式(10)所示.

$$Y_h(j+1)=\frac{Y_h^i[r_h\times L_h(j)-1]}{r_h\times L_h(j)+1}+\frac{r_h\times L_h(j)\times c}{r_h\times L_h(j)+1}\times$$
$$\sum_{k\neq h}^k c\frac{(B_{ui}-B_{li})f(p_k^i-p_h^i)(p_k-p_h)}{2D_{h,k}}$$

(10)

使用设计的混合优化算法求解最优解, 直到达到最大迭代限制为止, 输出最优权值, 将最优权重用于训练深度学习模型.

本研究所提方法如图 3 所示, 双模深度学习模型中 DBN 的权值有 w , 对于最优权值的获取, 将深度学习模型的权值归结为蝗虫位置. 首先初始化混合优化算法的参数, 初始化蝗虫种群位置, 然后再计算个体蝗虫的适应度值, 更新迭代次数, 如果未达到最大迭代次数, 计算个体之间的距离及单位距离向量, 将 CSA 作为种群中优势引导个体的趋势, 调整算法全局搜索和局部寻优能力. 如果适应度函数没有达到全局最优, 会产生新的位置, 如果新位置可行, 按照式(10)进行位置更新, 之后计算新位置的适应度函数, 直至达到最大迭代次数, 然后输出最优解, 并将其值赋给权值, 即可得到所需权值.

在求解最优权值过程中, 全局最优的具体位置并不知道, 将 CSA 的搜索算法引入到 GOA, 作为前进趋势, 使得蝗虫向着最优解的方向移动, 最终得到全局最优解, 即深度学习模型的最优权值.

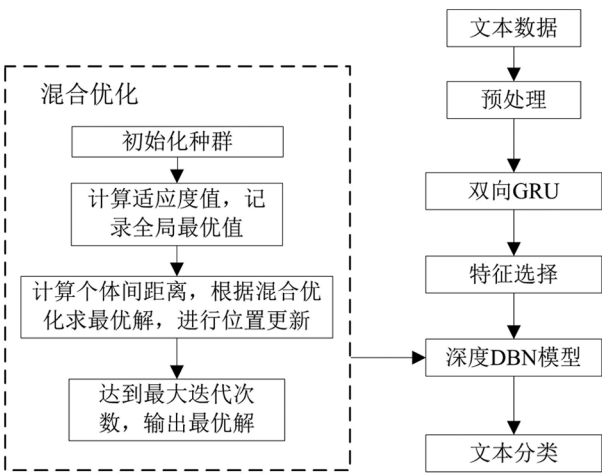


图 3 基于混合优化的双模深度学习文本分类

3 结果与分析

在 8GB RAM, Intel i7 处理器和 Windows 10 的计算机上进行试验. 性能指标有: 准确率、精确率、召回率.

$$\text{准确率}=\frac{T_p+T_n}{T_p+T_n+F_p+F_n}$$

(11)

$$\text{精确率}=\frac{T_p}{T_p+F_p}$$

(12)

$$\text{召回率} = \frac{T_p}{T_p + F_n}$$

(13)

其中， T_p 表示真阳性， T_n 表示真阴性， F_p 表示假阳性， F_n 表示假阴性。

为了验证所提混合优化算法的性能，将本文方法与传统乌鸦搜索 CSA 算法进行对比。使用包含单峰、多峰等不同特征的 6 种测试函数对混合优化算法和传统 CSA 算法进行试验，6 种测试函数具体设置如表 1 所示。为降低试验过程中因偶然性带来的误差，将本文混合算法与 CSA 算法对 6 种测试函数分别进行 50 次试验，并取其最差值、最优值、平均值及标准偏差，试验结果如表 2 所示。

表 1 不同测试函数

对应标号	函数名称	范围	最优值
f_1	Ackely 函数	$[-32, 32]$	0
f_2	Rastrigin 函数	$[-5.12, 5.12]$	0
f_3	Griewangk 函数	$[-600, 600]$	0
f_4	Sphere 函数	$[-100, 100]$	0
f_5	Schwefel's problem 22 函数	$[-10, 10]$	0
f_6	Sum Squares 函数	$[-10, 10]$	0

表 2 不同测试函数的试验结果

函数	项目	CSA	本文
f_1	最优值	2.670	8.88×10^{-16}
	最差值	4.780	8.88×10^{-16}
	平均值	3.610	8.88×10^{-16}
	标准差	0.638	0
f_2	最优值	7.990	0
	最差值	39.800	0
	平均值	23.900	0
	标准差	8.030	0
f_3	最优值	2.740×10^{-3}	0
	最差值	2.850×10^{-2}	0
	平均值	1.310×10^{-2}	0
	标准差	7.800×10^{-3}	0
f_4	最优值	0.027 1	5.21×10^{-102}
	最差值	0.224 0	9.64×10^{-84}
	平均值	0.078 6	5.39×10^{-85}
	标准差	0.047 8	2.08×10^{-84}
f_5	最优值	0.649	1.55×10^{-48}
	最差值	3.790	9.51×10^{-42}
	平均值	1.660	5.12×10^{-43}
	标准差	0.861	2.07×10^{-42}
f_6	最优值	0.188	3.91×10^{-102}
	最差值	3.370	1.61×10^{-84}
	平均值	1.100	8.07×10^{-86}
	标准差	0.958	2.52×10^{-85}

从表中数据可以看出,本文提出的混合优化算法对 6 种测试函数都有较好的最优值、最差值、平均值以及标准差,能够提供更好的收敛速度.这是因为本文混合优化算法在 CSA 中引入 GOA,解决了未知搜索空间的实际问题,在提供更好的收敛速度的同时获得全局最佳解决方案.

本研究验证本文混合优化算法的优越性,将本文算法与其他元启发式优化算进行比较,对比萤火虫、布谷鸟、蝙蝠、和声搜索和乌鸦搜索算法,测试函数选用 f_5 . 对比结果见表 3.

表 3 不同优化算法的性能对比

优化算法	平均值	优化算法	平均值
萤火虫	9.59×10^{-6}	和声搜索	9.69×10^4
布谷鸟	5.54×10^{-3}	乌鸦搜索 CSA	1.66×14^0
蝙蝠	2.68×10^1	本文	5.12×10^{-43}

从表中数据可以看出,本文算法具有最好的平均值,在提供更好收敛速度的同时获得全局最佳解决方案,说明本研究中的混合优化算法具有优越性.

试验数据集有两个,分别是 20 newsgroups 数据集和复旦大学中文文本分类库,20 newsgroups 数据集包括了 20 种新闻文本的 2 万篇新闻稿;复旦大学中文文本分类库也具有 20 个类的文档,文档类别包括农业、军事、电脑、艺术、运动、经济等各种领域.图 4 和图 5 分别给出了不同数据集中文本分类的准确率.

图 4 中 A~T 分别表示 20 newsgroups 数据集的 20 种类:alt.atheism, comp.graphics, comp.os.ms-windows.misc, comp.sys.ibm.pc.hardware, comp.sys.mac.hardware, comp.windows.x, misc.forsale, rec.autos, rec.motorcycles, rec.sport.baseball, rec.sport.hockey, sci.crypt, sci.electronics, sci.med, sci.space, soc.religion.christian, talk.politics.guns, talk.politics.mideast, talk.politics.misc, talk.religion.misc.

从图 5 中可以看出本文算法的分类准确率整体较高,英文分类准确率达到 85%以上,中文分类准确率达到 83%以上,表明算法的有效性.另外,英文分类准确率高于中文分类准确率,这是因为中文文本分类涉及中文分词问题,增加了分类复杂性,导致分类准确率降低.

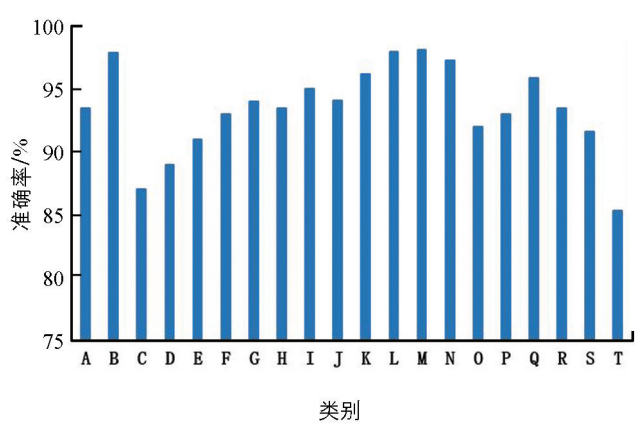


图 4 英文文本分类的准确率

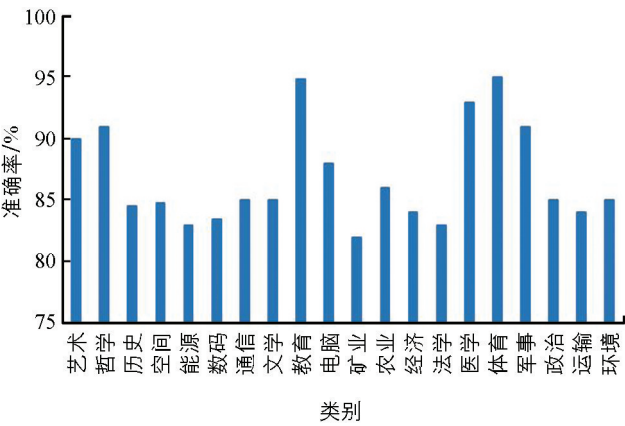


图 5 中文文本分类的准确率

为了验证本文方法的先进性,将本文方法与现有其他方法进行比较,对比方法有:文献[7]中改进卡方分布的文本分类方法,文献[13]中卷积递归神经网络的文本分类方法,文献[14]中胶囊网络和 LSTM 的文本分类方法.针对英文文本数据集,图 6 至图 8 给出了不同方的性能对比.

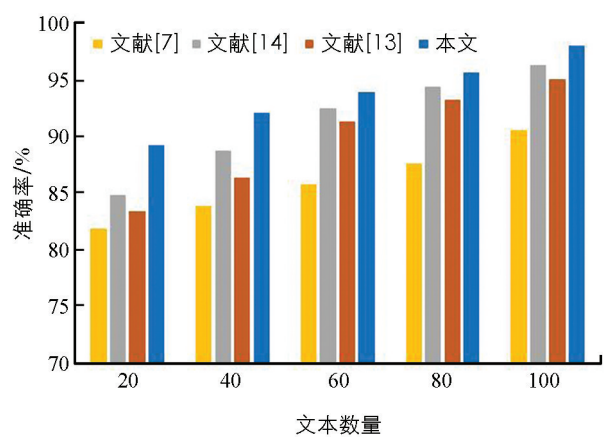


图 6 不同方法的文档分类准确率

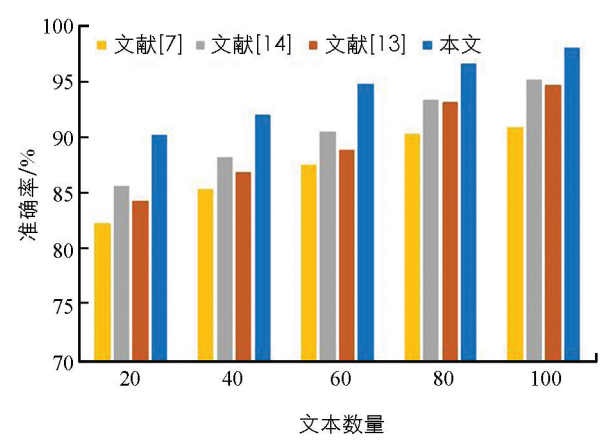


图 7 不同方法的文档分类精确率

研究结果可以看出，本文方法在准确率、精确率和召回率性能方面优于现有其他方法，这是因为本文方法在特征选择上使用了双向 GRU 深度模型，得到相关度较高的特征，然后使用混合优化的 DBN 进行文本特征的二次分析，有效地减少了分类类别，提高分类精度。文献[7]中基于改进卡方的文本特征分类方法性能最差，这是因为传统的卡方分布方法得到的特征相关度低；文献[13]中运用基于卷积递归神经网络的文本分类，使用卷积神经网络进行特征提取，在处理更长且复杂的文本结构和文本语句时，会因为卷积核长度的限制对相关特征提取造成影响，另外卷积神经网络进行池化时会丢失空间信息，这些都会影响文本分类精度；而文献[14]中胶囊网络和 LSTM 的文本分类方法，虽然使用了双模的神经网络模型进行文本分类，但是该方法中对深度模型进行参数寻优，使得分类性能仅次于本研究方法。

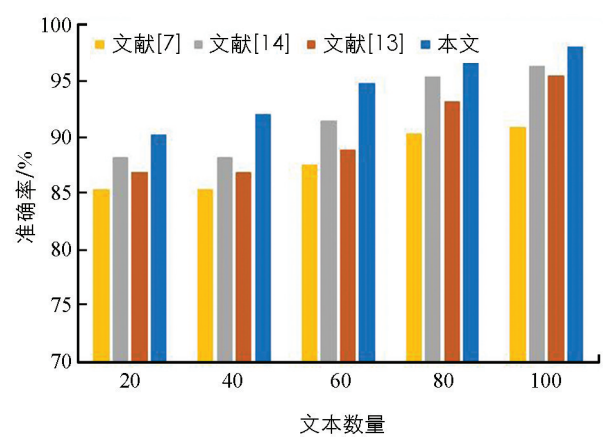


图 8 不同方法的文档分类召回率

4 结论

本研究提出了一种混合优化的双模深度学习文本分类方法。首先使用停用词删除和词干提取技术对输入文档进行预处理，以使输入有效并具有特征提取过程的能力；设计了基于 GOA 和 CSA 的混合优化算法，使用双向 GRU 深度学习模型进行特征选择；将选择的特征输入到混合优化的 DBN 深度模型中进行文本分类。试验结果表明，本研究中混合优化算法的双模深度学习文本分类方法能够以较高的性能实现文本分类，且性能优于现有方法。未来工作中可对混合优化算法中的适应度函数进行研究，进一步提高文本分类性能。

参考文献：

[1] DLIGACH D, AFSHAR M, MILLER T. Toward a Clinical Text Encoder: Pretraining for Clinical Natural Language Processing with Applications to Substance Misuse [J]. Journal of the American Medical Informatics Association, 2019, 26(11): 1272-1278.

[2] ZULQARNAIN M, GHAZALI R, GHHOUSE M G, et al. Efficient Processing of GRU Based on Word Embedding for

- Text Classification [J]. JOIV: International Journal on Informatics Visualization, 2019, 3(4): 377-383.
- [3] 唐焕玲, 窦全胜, 于立萍, 等. 有监督主题模型的 SLDA-TC 文本分类新方法 [J]. 电子学报, 2019, 47(6): 1300-1308.
- [4] KOU G, YANG P, PENG Y, et al. Evaluation of Feature Selection Methods for Text Classification with Small Datasets Using Multiple Criteria Decision-Making Methods [J]. Applied Soft Computing, 2020, 86: 105-116.
- [5] DENG X L, LI Y Q, WENG J, et al. Feature Selection for Text Classification: a Review [J]. Multimedia Tools and Applications, 2019, 78(3): 3797-3816.
- [6] 赵婧, 邵雄凯, 刘建舟, 等. 文本分类中一种特征选择方法研究 [J]. 计算机应用研究, 2019, 36(8): 2261-2265.
- [7] BAHASSINE S, MADANI A, AL-SAREM M, et al. Feature Selection Using an Improved Chi-Square for Arabic Text Classification [J]. Journal of King Saud University - Computer and Information Sciences, 2020, 32(2): 225-231.
- [8] JOSEPH MANOJ R, ANTO PRAVEENA M D, VIJAYAKUMAR K. An ACO-ANN Based Feature Selection Algorithm for Big Data [J]. Cluster Computing, 2019, 22(2): 3953-3960.
- [9] KOWSARI K, MEIMANDI J, HEIDARYSAFA M, et al. Text Classification Algorithms: a Survey [J]. Information, 2019, 10(4): 150-217.
- [10] 韩众和, 夏战国, 杨婷. CNN-ELM 混合短文本分类模型 [J]. 计算机应用研究, 2019, 36(3): 663-667, 672.
- [11] SINGH G, KUMAR B, GAUR L, et al. Comparison between Multinomial and Bernoulli Naïve Bayes for Text Classification [C] //2019 International Conference on Automation, Computational and Technology Management (ICACTM). London, UK. IEEE, 2019: 593-596.
- [12] LIU G, GUO J B. Bidirectional LSTM with Attention Mechanism and Convolutional Layer for Text Classification [J]. Neurocomputing, 2019, 337: 325-338.
- [13] WANG R S, LI Z, CAO J, et al. Convolutional Recurrent Neural Networks for Text Classification [C] //2019 International Joint Conference on Neural Networks (IJCNN). Budapest, Hungary. IEEE, 2019: 1-6.
- [14] 胡春涛, 夏玲玲, 张亮, 等. 基于胶囊网络和卷积网络的文本分类对比 [J]. 计算机技术与发展, 2020, 30(10): 86-91.
- [15] 后同佳, 周良. 基于双向 GRU 神经网络和注意力机制的中文船舶故障关系抽取方法 [J]. 计算机科学, 2021, 48(S2): 154-158.
- [16] HU C H, PEI H, SI X S, et al. A Prognostic Model Based on DBN and Diffusion Process for Degrading Bearing [J]. IEEE Transactions on Industrial Electronics, 2020, 67(10): 8767-8777.

责任编辑 王新娟