

# 基于卷积神经网络的数字化教育评论情感分类应用

田茗瑶, 陈 洋

(河北工程大学数理科学与工程学院, 河北 邯郸 056038)

**摘要:** 随着科技发展和数字化转型, 如何理解和提升学生的情感体验是教育工作者和政策制定者关注的重点。首先对 2 000 条人工标注的评论数据使用三种机器学习模型进行训练, 利用已训练支持向量机(SVM)模型对 76 000 条评论数据进行自动标注, 并使用这些标注数据集训练卷积神经网络(CNN)模型。CNN 模型在 10 次迭代后成功收敛, 并在验证集上取得 94.96% 的准确率, 显著优于 SVM 模型。结果显示, 结合传统机器学习与深度学习的方法可有效提高情感分析的准确性, 为教育策略优化提供数据支持。

**关键词:** 卷积神经网络; 数字化教育; 情感分析; 机器学习

**中图分类号:** TP391.1; TP183; G40-01 **文献标志码:** A **文章编号:** 1671-1807(2025)02-0121-07

随着科技的迅速发展和数字化转型的推进, 数字化学习逐渐成为现代教育体系的重要组成部分。线上教育平台的普及, 极大地改变了传统的教学方式, 也带来了丰富的学生反馈数据<sup>[1-2]</sup>。这些反馈数据, 特别是学生对线上课程的评论和意见, 蕴含着大量关于课程质量、教学效果以及学生学习体验的情感信息。如何高效、准确地分析这些情感信息, 已成为教育工作者和政策制定者关注的焦点。

情感分类作为自然语言处理的重要分支, 旨在从文本数据中自动提取用户的情感倾向。传统的情感分类方法主要分为基于词典的方法和基于机器学习的方法<sup>[3]</sup>。基于词典的方法依赖于预定义的情感词典, 通过匹配文本中的情感词来判断情感倾向; 而基于机器学习的方法, 使用人工标注的数据, 通过支持向量机(support vector machine, SVM)、朴素贝叶斯等算法对文本进行分类。然而, 这些方法在处理大规模复杂数据时存在显著的局限性, 尤其是大量人工标注的需求不仅耗费时间和成本, 且标注效果难以保证一致性<sup>[4,5]</sup>。

近年来, 深度学习技术的迅猛发展为情感分析带来了新的契机。卷积神经网络(convolutional neural networks, CNN)作为深度学习的重要模型之一, 凭借其在图像处理和特征提取方面的优越性能, 逐步被应用于文本情感分类领域。CNN 能够自动从数据中学习出高层次的特征, 在处理大规模数据时展现出卓越的性能。

为应对大量人工标注的挑战, 本文首先选择了适合处理小规模数据集的传统机器学习模型, 对 2 000 条人工标注的评论数据进行训练。通过这种方法, 能够减少人工标注的负担, 同时确保标注的准确性和一致性。在此基础上, 利用经过训练的 SVM 模型对大规模评论数据进行自动标注。

之后, 采用 CNN 模型对这些自动标注的数据进行情感分类。CNN 在特征提取和处理大规模数据集方面表现出色, 显著提升了情感分类的整体性能。实验结果表明, CNN 模型在验证集上取得了更高的准确率, 且在大规模数据集上的表现优于传统的机器学习方法。

本文通过结合传统机器学习和深度学习的方法, 不仅提升了情感分析的准确性, 也有效减少了人工标注的工作量, 为线上教育评论的情感分析提供了一种高效且可靠的解决方案。

## 1 数据获取及预处理

### 1.1 数据来源

通过对中国五大社交媒体平台——Bilibili(b 站)、抖音、快手、小红书和微博的评论进行情感分析, 深入了解学生对线上学习的看法和感受。这些平台覆盖了广泛的用户群体, 从小学生到大学生, 均在这些平台上积极发表关于线上教育的评论和反馈。

数据采集利用 Python 编写的网络爬虫程序, 模拟用户登录, 并通过平台的搜索接口根据预设关键词(如“线上学习”“远程教育”等)捕捉相关的评论。

数据收集时间跨度为 2020—2024 年, 共计爬取

**收稿日期:** 2024-08-25

**作者简介:** 田茗瑶(2002—), 女, 河北正定人, 研究方向为自然语言处理与情感分析; 陈洋(2004—), 男, 江苏扬州人, 研究方向为自然语言处理与情感分析。

了约 76 000 条有效评论。这些评论不仅提供了学生对线上学习模式的直接反馈,也反映了他们对于教育内容、教学质量及互动体验的看法。

通过利用社交媒体平台作为数据来源,获取了丰富、多维度的实时数据资源,有助于全面理解学生对数字化教育的态度和需求。这些洞察对于制定教育政策和优化教育模式提供了有力的支持,从而更好地适应当代教育环境的变化。

部分数据集如图 1 所示。

1.2 数据清洗及预处理

在进行数据分析之前,首先需要对非结构化的文本数据进行预处理,以转换为计算机可处理的结构化信息。此阶段的关键步骤包括分词和去除停用词,这些操作有助于清理和规范化数据,为之后的深入分析打下坚实的基础。

1.2.1 分词

在进行中文文本处理前,由于中文缺乏明确的词界标识,如空格等,使得词与词之间的界限不明显。因此,分词成为文本分析中一个关键且必不可少的步骤。准确的分词不仅有助于提升后续文本挖掘的质量,还是获取有效信息的基础。

本文采用 Python 的 jieba 库进行中文分词,以有效处理并分析从社交媒体平台收集到的评论数据。选择精确模式进行分词,这种模式能够精确地将句子切分成词语,避免冗余和歧义的产生,非常适合用于精细的文本分析(表 1)。

表 1 Jieba 分词示例

|     |                  |
|-----|------------------|
| 原句  | 我们要开线上课程啦        |
| 分词后 | 我们 要 开 线 上 课 程 啦 |

| 标题                | 评论             | 发布日期                | 用户id       | 用户名         | 点赞数               | 转发数   | 评论数   |
|-------------------|----------------|---------------------|------------|-------------|-------------------|-------|-------|
| 【大学网课】我在b站上       | 瑟瑟发抖,在零度的走     | 2023-11-12 14:19:21 | 23310707   | Donut甜甜圈    | 18 334            | 335   | 371   |
| 帮你省钱 我做了7年在       | 关注我私信关键词【学     | 2024-04-02 06:10:28 | 39460908   | 问问大象        | 8 863             | 175   | 175   |
| 24考研数一135不听网课     | 有什么想问的再私集我     | 2024-04-22 15:17:03 | 32269835   | 禾沐呀         | 245               | 12    | 65    |
| 让我效率翻倍的听网课        | 给大家承诺过的干贷分     | 2021-03-26 08:45:17 | 495713183  | -炒肉丝-       | 5 485             | 302   | 238   |
| 【全748集】花了26 800买  | 内容都是在不耽误主要     | 2024-03-10 06:00:39 | 3546636672 | 团子学姐吖       | 715               | 101   | 21    |
| 因为网课,我的高三报废       | 不要像我这样!!不要     | 2022-01-22 12:08:59 | 373103803  | 树林同学-       | 2×10 <sup>5</sup> | 7 125 | 7 066 |
| 高中数学网课大合集         | 资源来源于网上,侵权     | 2021-02-07 05:16:13 | 1541017767 | 若不趁起风时      | 452               | 170   | 70    |
| 考研听盗版网课也行,记得资料找全点 |                | 2023-09-06 03:06:14 | 3493296991 | 大雁考研英语      | 2 107             | 1     | 143   |
| 只听网课,不听学校的可       | BV1H34y1R7Qm   | 2023-07-09 05:27:09 | 1043513102 | INSame_Zeus | 12 040            | 150   | 618   |
| 【2024年全100集】雅思    | 雅思口语的神Keith网   | 2024-04-14 06:00:21 | 502547872  | 小熊i顶呱呱      | 670               | 118   | 12    |
| 【雅思写作】2024年顾家     | 222视频上传不易,多    | 2024-04-25 12:09:36 | 346571327  | 重生之我在B站     | 9                 | 11    | 8     |
| 【大学网课】我在b站上       | 好久不见ww,前段时     | 2024-01-12 12:12:52 | 23310707   | Donut甜甜圈    | 2 070             | 173   | 81    |
| 【最新Simon听说读写全     | 2024前雅思考官Simon | 2024-04-16 11:19:20 | 3546389726 | 雅思小熊        | 541               | 124   | 106   |

图 1 部分数据集

1.2.2 去除停用词

经过分词处理后,文本数据中往往含有大量的停用词,如“的”“这个”“那个”等,这些词虽频繁出现,但对于深入分析文本并无实质帮助。因此为提升数据分析的质量和效率,需在词频统计前进行一项关键的数据清洗工作——去除停用词。

为了高效地执行这一步骤,首先构建了一个停用词列表。该列表融合了多个来源的停用词表,并针对社交媒体文本的特点进行了定制和优化。

清洗过程中,每一条经过分词的文本数据都将与停用词表列表进行匹配检查,所有列表中的词汇将被从文本中删除。这一步骤不仅清理了无关信息,也为提取每条评论中的关键词和进行深入分析奠定了坚实基础。

通过这种方法,确保了数据的清洁度和分析的准确性,为后续的文本分析和信息提取提供了高质量的输入。

分词及去除停用词后的部分结果如图 2 所示。

2 基于机器学习的情感分类

首先选择部分爬取的相关文本数据,进行细致的人工情感标注(1 表示积极,0 表示中性,-1 表示消极情绪),并创建了包含 2 000 条评论的专门语料库。这一数据规模既避免了过大数据量导致的模型训练复杂度增加和过拟合问题,同时也足够保证数据的多样性和全面性,从而有效地支持模型的训练和验证。通过精心设计的数据集,本文能够有效平衡模型的泛化能力和识别精度,适应教育环境下复杂多变的情绪分类需求。

部分语料库如图 3 所示。

|                                                                                                 |
|-------------------------------------------------------------------------------------------------|
| 帮省钱 做年在线教育没花钱买网课……                                                                              |
| 24 考研 数一 135 听网课野路子                                                                             |
| 全 748 集 花 26 800 买 PTE 网课 2024 最新 内部 版 包含 PTE 听说读写 全套 学完 即 八 炸 上岸 拿走 不 谢 学 退出 教学 圈              |
| 2024 年 全 100 集 雅思 口语 神 Keith 网课 中文字幕 版 含 配套 讲义 零 基础 雅思 小白 必 油管 最新 雅思 口语 课程                      |
| 雅思 写作 2024 年 顾家 北 手把手 教 雅思 写作 网 课 提 分 必 看 附 带 PDF 写 作 范 文 观 点 库                                 |
| 大学 网 课 b 站 上 大学 → 一 宝 藏 网 课 神 仙                                                                 |
| 最新 Simon 听说读写 全集 2024 前 雅思 考官 Simon 雅思 中文字幕 版 网 课 附 讲 义 Simon 写 作 Simon 口 语 Simon 听 力 Simon 阅 读 |
| B 站 独 家 系 统 公 务 员 网 课 告 别 低 劣 教 程 零 基 础 带 少 走 99% 弯 路 学 滚 出 考 公 圈                               |
| B 站 考 研 英 语 网 课 性 价 比 之 王                                                                       |
| 2024 年 顾 家 北 雅 思 写 作 网 课 全 集 手 把 手 教 提 分 必 看 附 pdf 满 分 范 文                                      |
| 上海 体育 大学 网 球 选 修 课 学 生 有 多 强 30 节 网 课 后 现 状                                                     |
| 雅思 冒 死 上 传 N 次 名 师 顾 家 北 24 年 最 新 最 全 雅 思 写 作 网 课 合 集 手 把 手 教 雅 思 写 作 附 高 清 网 课 配 套 讲 义         |
| 考 研 英 语 网 课 性 价 比 之 王                                                                           |
| 北 师 大 老 师 发 传 单 宣 传 网 课                                                                         |
| 高 考 网 课 480 暴 涨 642 分 逆 袭 211 网 课 名 师 推 荐 网 课 避 雷 网 课 双 面 性                                     |
| 上 海 管 综 250 管 综 数 学 看 网 课 效 率 网 课                                                               |
| 高 中 生 听 网 课 越 听 分 越 低                                                                           |
| 雅 思 网 课 冒 死 上 传 52 次 离 职 堪 称 B 站 完 整 雅 思 网 课 包 含 听 说 读 写 没 人 更                                  |
| bilibili 大 学 最 全 雅 思 网 课 闭 眼! 内 容 全 程 干 货 雅 思 听 说 读 写 全 科 魔 鬼 训 练 附 配 套 资 料                     |
| 25 考 研 数 学 网 课 网 盘 免 费 分 享 课 最 新                                                                |

图 2 分词及去除停用词结果

|    |                                                   |
|----|---------------------------------------------------|
| 1  | 高考网课评分第二弹强势来袭，比以往更加客观                             |
| 1  | 2345祝大家早日雅思上岸                                     |
| -1 | 睡觉，刷抖音，聊天，照镜子，一提问啥都不会，上网课的我的真实状态                  |
| 0  | 小蹭一下沙雕之主的热度，不想当rapper的教书匠不是好生物老师                  |
| 1  | B站的优秀老师推荐，本人非营销号，一个高考后做自媒体的学生而已大家偷偷收藏内卷你们的同学这可    |
| -1 | #现实版搞笑一家人#女儿在做网课老师布置的视频作业，爸爸以为她在录着玩，所以突然在背景里舞了起来  |
| 0  | 4月7日，同济大学数学科学学院一则上网课的视频引发热议。据了解，当时教授李雨生正在给学生们上网课  |
| -1 | 本期1，到底要不要做老师发的模拟题2，没时间看网课怎么办3，想提升XX分有没有希望         |
| -1 | 唠点肺腑之言，只剪掉了一点说错话的地方，没有写稿子，纯唠嗑，爱看不看。从3分46开始是将给高中生的 |
| 1  | 本地上传好好学习，天天向上，早日雅思上岸                              |
| -1 | 天天上网课太无聊?教你如何科学地开小差! (误)                          |
| 1  | 高中网课，老师推荐合集                                       |

图 3 部分语料库

采用三种不同的机器学习算法——朴素贝叶斯、支持向量机(SVM)及随机森林来构建情感分类模型。算法的选择是在广泛查阅相关文献后进行的,这些算法被广泛认为在文本情感分类任务中表现良好,尤其是在处理高维数据时,其在各自的领域内已经被证明能有效地提高情感分类的准确性和可靠性,因此将其运用在自建的语料库上进行测试和比较<sup>[6-7]</sup>。

2.1 朴素贝叶斯

朴素贝叶斯的核心思想在于利用特征之间相互独立的假设<sup>[6]</sup>。在训练阶段,计算先验概率、条件概率和联合概率分布。在预测阶段,通过应用贝叶

斯定理来计算后验概率,并根据后验概率的大小确定测试数据的分类结果。其训练和分类过程大致如下:

(1)对于任意一个数据样本,  $n$  维特征向量  $\mathbf{X}(x_1, x_2, \dots, x_n)$  表示  $n$  个特征属性的特征向量。输出类别分别为  $C_1, C_2, \dots, C_m$ , 共计  $m$  个类别。

(2)计算类先验概率:对于类别  $C_i(1 \leq i \leq m)$ , 先验概率的公式为

$$P(C_i) = \frac{N_i}{N} \tag{1}$$

式中: $N$  为训练集中样本总数。

(3)计算条件概率:朴素贝叶斯算法是基于条

件独立的基础,即

$$P(\mathbf{X} | C_i) = P(x_1, x_2, \dots, x_n | C_i) = P(x_i | C_i)P(x_2 | C_i) \cdots P(x_n | C_i) = \prod_{k=1}^n P(x_k | C_i) \quad (2)$$

式中:  $k$  用于标记第  $k$  个特征值  $x_k$ ;  $P(x_i | C_i)$  ( $1 \leq i \leq k, 1 \leq i \leq m$ ) 由训练过的数据计算得出。条件概率计算公式为

$$P(x_k | C_i) = \frac{N(x_k, C = C_i)}{N(C = C_i)} \quad (3)$$

(4) 计算后验概率: 基于贝叶斯定理, 可推出后验概率  $P(C_i | \mathbf{X})$  的计算公式为

$$P(C_i | \mathbf{X}) = \frac{P(\mathbf{X} | C_i)P(C_i)}{P(\mathbf{X})} \quad (4)$$

式中:  $P(\mathbf{X} | C_i)$  为条件概率;  $P(C_i)$  为类先验概率;  $P(\mathbf{X})$  为常数。

(5) 输出类别: 比较全部后验概率的大小, 朴素贝叶斯分类器对未分类的样本进行分类, 最终输出结果是类别  $C_i$ , 当且仅当  $C_i$  满足:  $P(C_i | \mathbf{X}) < P(C_j | \mathbf{X}), 1 \leq i \leq m, j \neq i$ 。

预处理阶段:

该阶段是由人工完成的, 需要手动定义特征属性, 然后适当地对特征进行类别划分, 以形成一个训练样本集(部分分类的特分类项集合)一个分类器的好坏很大程度上也取决于训练样本集的质量

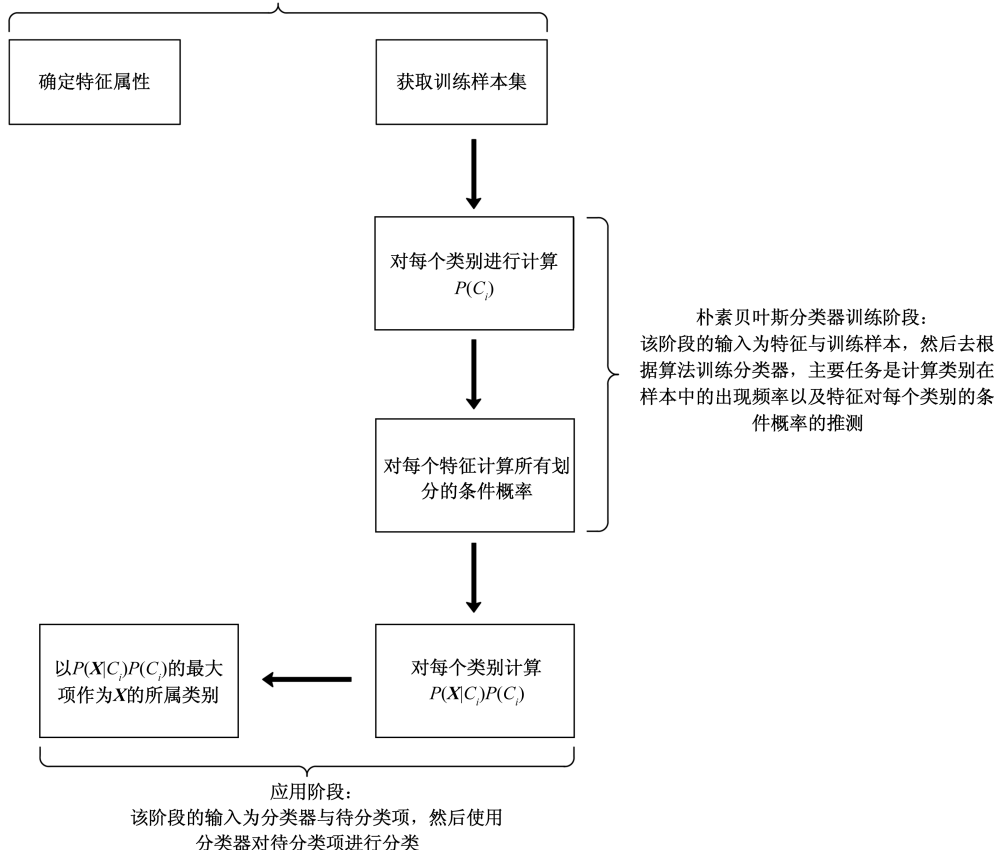


图 4 朴素贝叶斯算法流程

具体流程如图 4 所示。

首先从“自建语料库.xlsx”文件中加载数据。为了保证数据质量, 在加载数据后首先对其进行清洗, 删除包含空值的行。此外, 为了提高文本数据的处理效率和效果, 使用 jieba 分词库对文本内容进行分词处理。分词后的文本以空格分隔单词, 便于后续的向量化处理。

在文本预处理完成后, 使用 TfidfVectorizer 对文本数据进行向量化。TF-IDF 向量化不仅考虑词频(term frequency, TF), 也考虑词语在整个数据集中的逆文档频率(inverse document frequency, IDF), 这有助于减少常见词的影响, 同时突出重要词汇。通过这种方式, 每条文本数据被转换为一个数值型特征向量, 为模型训练提供了适合的输入格式。

采用朴素贝叶斯分类器(MultinomialNB)进行模型训练。选择  $\alpha=0.5$  作为平滑参数, 该参数有助于调节模型对未见词的处理, 从而避免过拟合。数据集被分为 70% 的训练集和 30% 的测试集, 使用随



机种子 42 来保证结果的可重复性。模型训练完成后,对测试集进行预测,并通过 classification\_report 输出了模型的性能评估,包括准确率、召回率和 F1 分数等指标,这些指标有助于理解模型在各类情感分类任务上的表现。

经过训练,模型在所有情绪类别上的准确率为 86%。加权平均的准确率、召回率和 F1 分数均在 83%至 86%之间,显示出模型在处理不同情绪类别时的均衡性和有效性。

## 2.2 支持向量机

支持向量机(SVM)是一种常见的二分类算法,可用于解决二分类和多分类问题,SVM 在文本分类任务中表现出色,因其可有效地处理高维数据和小样本问题。继朴素贝叶斯方法之后,本文进一步采用支持向量机(SVM)进行文本情感分类。

在实施支持向量机之前,利用同一 TF-IDF 向量化的文本数据作为输入。选择线性核函数的 SVM 模型,因为线性模型在文本分类任务中通常能提供良好的结果,同时计算效率较高。

设置正则化参数  $C=1.0$ ,该设置为经验值,通常能够在训练和泛化之间达到良好的平衡。模型训练采用了与朴素贝叶斯相同的数据划分,即 70%的数据用于训练,30%用于测试。

训练完成后,支持向量机的性能通过相同的性能评估指标进行了评估。结果表明,SVM 在总体准确率上达到了 91%,且在召回率和 F1 分数方面也表现出较高的水平。

## 2.3 随机森林

在朴素贝叶斯和支持向量机(SVM)的研究基础上,进一步采用随机森林算法来进行文本情感分类,以探索和比较不同机器学习算法在同一数据集上的性能表现。

随机森林是一种集成学习技术,通过构建多个决策树并汇总它们的预测结果来提高整体性能,特别是在减少过拟合和提高预测准确性方面表现出色。

首先配置了一个由 100 棵决策树组成的随机森林模型。这些决策树在训练过程中随机选择特征子集进行学习,这不仅增强了模型的泛化能力,也有效避免了单一决策树可能出现的过拟合问题。与前两种模型一样,随机森林也利用了 TF-IDF 向量化的文本数据作为输入,确保了实验的一致性。数据的分割策略维持为 70%用于训练,30%用于测试。

通过运行,随机森林模型在所有情绪类别上的整体准确率达到 90%,表现出较好的平衡性和有效性。加权平均的准确率、召回率和 F1 分数均为 90%~91%,显示出该模型在处理不同情绪类别时的一致性和效率。

## 2.4 基于不同分类模型的实验结果分析

通过使用三种不同的机器学习算法——朴素贝叶斯、支持向量机(SVM)和随机森林,旨在深入探索这些算法在该情感分类任务上的表现。每种算法都有其独特的处理方法和对高维数据的不同适应性,以期找到最适合分析和解读数字化教育环境中学生情绪反应的技术<sup>[5-6]</sup>。

特别关注这些模型在准确性、计算效率和对不平衡数据的处理能力方面的表现,旨在识别出哪一种模型不仅能提供高精度的情感识别,同时也能适应教育领域数据分析的实际需求。通过对比这些模型在同一数据集上的性能,希望揭示各算法的优势和局限,为未来在数字化教育平台上实施情感分析提供科学依据和方法指导。

表 2 展示了朴素贝叶斯、支持向量机(SVM)和随机森林这三种机器学习算法在本次情感分类任务中的性能对比,包括各自的准确率、召回率和 F1 分数,从而提供了一个直观的性能评估视角。

表 2 不同分类模型的实验结果

| 分类器   | 准确率  | 召回率  | F1 分数 |
|-------|------|------|-------|
| 朴素贝叶斯 | 0.86 | 0.85 | 0.83  |
| 支持向量机 | 0.91 | 0.91 | 0.91  |
| 随机森林  | 0.90 | 0.91 | 0.90  |

在本次情感分类任务中,通过比较朴素贝叶斯、支持向量机(SVM)、和随机森林三种不同的机器学习算法,分析结果表明,支持向量机在所有性能指标上均表现最优,整体准确率、召回率和 F1 分数均为 0.91,显示其在高维数据处理和复杂模式识别中的优势,但训练速度较慢,对参数调整较为敏感。朴素贝叶斯以 0.86 的准确率较快地完成任务,适用于初步分析,但在处理相关特征时效果受限。随机森林的准确率达 0.90,表现出良好的鲁棒性和对大规模数据的适应性,不过计算资源消耗相对较大。

尽管每种模型都有其优点和局限,支持向量机(SVM)在本文表现最为突出,显示了其在处理数字化教育环境下学生情感态度分类时的高准确性和良好的平衡性。

3 基于卷积神经网络的情感分类模型

3.1 情感标注

基于上述工作,选择已训练的 SVM 模型对全部爬取的 76 000 条评论进行自动情感标注。目的是利用支持向量机模型对大规模数据集进行快速且准确的标注,该过程显著减少了人工标注的工作量,并且借助已训练模型的高准确性,确保了大规模数据标注的可靠性,为后续的深度学习模型训练提供更为丰富的数据支持。

3.2 基于 CNN 模型训练

采用 PyTorch 框架构建了一个用于文本情感分类的卷积神经网络(CNN)模型<sup>[8-9]</sup>。模型包括输入层、嵌入层、卷积层、全局池化层、全连接层和输出层。嵌入层的维度设定为 64,用于捕捉输入文本的语义信息。卷积层使用了 100 个卷积核,每个卷积核的尺寸为 6,以便有效提取文本中的局部特征。卷积操作后的输出通过全局池化层进行降维,从而保留关键特征。随后,这些特征被送入由 100 个节点组成的全连接层,用于进一步提取高级抽象特征。输出层采用 softmax 激活函数,将每条文本归类为三类情感(积极、中性、消极)中的一种(图 5)。

为了优化模型参数,模型通过反向传播算法和 Adam 优化器进行训练。在此过程中,为了找到最优的超参数组合,使用网格搜索法通过调整学习率(learning rate)和批量大小(batch size)的不同组合,并以训练一轮在验证集上的准确率作为评估指标来选择最佳参数。

首先设置多个学习率的候选值(0.01、0.001、0.000 1),以及不同的批量大小(32、64、128)。对于每一个超参数组合,对模型进行训练,并记录每个组合在验证集上的表现。通过对所有组合的结

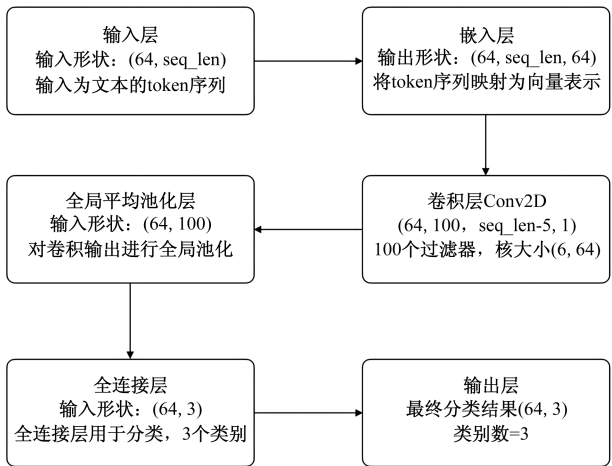


图 5 卷积神经网络模型流程

果进行对比分析,最终确定了学习率为 0.001 和批量大小为 64 的组合在验证集上取得了最高的准确率。

表 3 展示了不同学习率和批量大小组合下的验证集准确率,通过表 3 可以清晰地看到每个超参数组合的效果,并直观地找出最优的参数设置。

接着在整个训练过程中,模型的学习率设为 0.001,批量大小为 64。为了保证模型的泛化能力,整个数据集被划分为训练集和验证集。模型在 10 次迭代中训练,每次迭代时通过交叉熵损失函数对模型进行监督学习,并根据验证集的表现评估模型的泛化性能。最终,通过比较验证集上的准确率、精度、召回率等指标,验证了模型在情感分类任务中的有效性。训练过程中,模型的损失函数收敛情况和验证集的准确率如图 6 所示。

3.3 模型性能对比

我们对训练的神经网络模型(CNN)与传统机器学习模型支持向量机(SVM)在情感分类任务上的表现进行了对比分析。模型的性能通过准确率(accuracy)、召回率(recall)和  $F_1$  分数( $F_1$  Score)三个主要指标来评估。表 4 总结了这两种模型在验证集上的表现。

表 3 模型超参数组合的验证集准确率

| 批量大小 | 准确率/%       |            |           |
|------|-------------|------------|-----------|
|      | 学习率 = 0.001 | 学习率 = 0.01 | 学习率 = 0.1 |
| 32   | 91.55       | 91.11      | 91.45     |
| 64   | 91.65       | 90.89      | 90.57     |
| 128  | 91.14       | 91.26      | 91.51     |

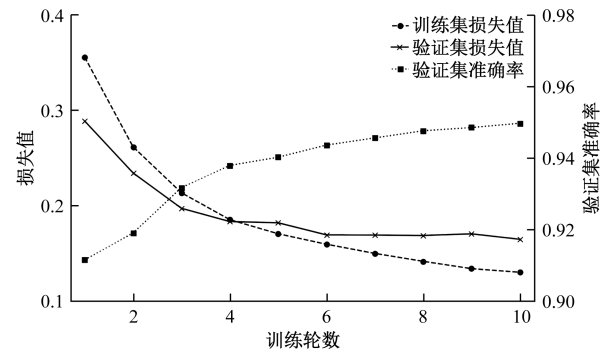


图 6 CNN 模型的准确率曲线和损失值曲线

表 4 模型性能对照

| 模型  | 准确率  | 召回率  | $F_1$ 分数 |
|-----|------|------|----------|
| SVM | 0.91 | 0.91 | 0.91     |
| CNN | 0.95 | 0.94 | 0.94     |

根据表 4 中对比的两种模型在情感分类任务上的性能表现可以发现,随着语料库的扩充和更深度的训练,模型的效果得到了显著提升。相较于传统的支持向量机(SVM)模型,卷积神经网络(CNN)在多个评价指标上展示了更为优越的性能。

CNN 模型准确率达到 0.95,而 SVM 模型的准确率为 0.91。这一差距表明,CNN 在更大规模的数据集上能够更有效地学习和捕捉评论中的情感特征。特别是在处理复杂语言模式和上下文信息时,CNN 的深度特征提取能力显得更加突出。通过多层次的卷积操作,CNN 能够从评论文本中提取更丰富的特征表示,从而显著提高了分类的准确性。

在召回率方面,CNN 模型的表现也优于 SVM,达到了 0.94,相较于 SVM 的 0.91,CNN 在检测多样化情感表达时具有更高的覆盖率,能够有效减少漏检的情况。CNN 的多层次网络结构使其能够捕捉情感分类中的细微差异,因此在处理具有高复杂度和多样性的情感数据时表现更为优越。

CNN 模型不仅在准确性和召回率上表现出色,而且在应对不平衡数据集的情感分类任务中更具鲁棒性和适应性。CNN 通过卷积层的多通道操作和池化层的特征筛选,能够更高效地学习并优化情感特征,使其在整体性能上超越传统机器学习方法。

研究结果显示,随着语料库的扩充和深度学习方法的应用,卷积神经网络模型在情感分类任务中的表现显著优于传统的 SVM 模型。这表明 CNN 模型在准确率、召回率和 F1 分数等关键指标上具有更高的有效性和稳定性,进一步证明了其在大规模、复杂数据环境下进行情感分析的潜力和优势。这为未来在数字化教育领域更广泛应用深度学习技术进行情感分析提供了重要的依据。

## 4 结论

卷积神经网络模型在 10 次迭代中逐渐收敛,并在验证集上表现出较强的泛化能力,最终准确率达到 94.96%,明显优于传统机器学习模型。相较于 SVM 模型,CNN 在验证集上表现出更高的准确率和更低的损失值,能够更好地捕捉评论文本中的情感特征,并且在大规模数据集上有更好的表现。

本文展示了如何有效结合传统机器学习与深度学习模型,逐步提升情感分析的准确性,为线上教育评论的情感分析提供了有效的工具和经验,对相关领域的研究者和教育政策制定者具有参考价值。

## 参考文献

- [1] 张琨. 互联网推动教育数字化转型的逻辑机理与现实路径[J]. 互联网周刊, 2024(7): 27-29.
- [2] HOI S C H, SAHOO D, LUJ, et al. Online learning: a comprehensive survey[J]. Neurocomputing, 2021, 459: 249-289.
- [3] 关志广, 程乔. 基于 NLP 的文本挖掘技术在提升电信客户满意度中的应用[J]. 无线互联科技, 2023, 20(5): 117-119.
- [4] 桑倩, 王生花. 基于文本挖掘技术的学生评语情感分析[J]. 信息与电脑(理论版), 2022, 34(2): 38-41.
- [5] 张财, 马自强, 闫博. 基于机器学习的政务微博情感分析模型设计[J/OL]. 计算机工程, 1-14. [2024-08-01]. <https://link.cnki.net/urlid/31.1289.TP.20240423.1801.006>.
- [6] 邓慈云, 余国清. 基于朴素贝叶斯的影评情感分析研究[J]. 智能计算机与应用, 2023, 13(2): 210-212.
- [7] 孙昊男. 机器学习下网络平台文本的中文情感分析[J]. 黄河水利职业技术学院学报, 2023, 35(4): 49-56.
- [8] 黄丹婷. 基于集成学习的微博评论情感分析[D]. 大连: 大连交通大学, 2023.
- [9] 郑锐斌, 贺丹, 王凯, 等. 深度学习技术在高校网络舆情分析中的应用[J]. 福建电脑, 2024, 40(5): 21-26.

## Application of Sentiment Classification in Digital Education Reviews Based on Convolutional Neural Networks

TIAN Mingyao, CHEN Yang

(School of Mathematics and Physics, Hebei University of Engineering, Handan 056038, Hebei, China)

**Abstract:** With the advancement of technology and digital transformation, understanding and enhancing students' emotional experiences are key concerns for educators and policymakers. In this study, 2 000 manually labeled review data were first used to train three machine learning models. The trained SVM model was then used to automatically label 76 000 review data. These labeled datasets were subsequently used to train a Convolutional Neural Network (CNN) model. The CNN model converged successfully after 10 iterations and achieved an accuracy of 94.96% on the validation set, significantly outperforming the SVM model. The results show that combining traditional machine learning with deep learning methods can effectively improve the accuracy of sentiment analysis. This provides data support for optimizing educational strategies.

**Keywords:** convolutional neural network; digital education; sentiment analysis; machine learning