

Application of Text Analysis in Emotional Expression in Contemporary Literature

Jingwen Zhang¹, 

¹ School of Literature and Education, Shaanxi Institute of International Trade and Commerce, Xi'an, Shaanxi, 712046, China

ABSTRACT

With the rapid development of technology and online social networking, the popularization of smartphones has promoted the research and development of sentiment analysis of contemporary literary texts. In this paper, the CBOW model based on Hierarchical Softmax algorithm is used to extract text sentiment features. The classification mechanism of sentiment lexicon, machine learning, and deep learning methods supported by sentiment features is discussed. According to the discussion results, a 5-layer sentiment analysis model based on CNN-BiLSTM-ATT is built based on text preprocessing, and the model design of different layering is proposed. Meanwhile, the analysis method of text themes is proposed based on LDA. In the long story dataset, the model recall rate of this paper is 83.91% and the precision rate is 83.86%, the values are higher than the other six models; the Macro-F1 mean value is 83.16%, which proves that the fused and improved CNN-BiLSTM-ATT model of this paper possesses excellent performance in the sentiment analysis task. In the short story dataset, the accuracy, precision and recall are not less than 98%, and the loss rate is the lowest 34.11%, which are lower than the other six models. The model in this paper can be applied to text analysis systems and has superiority in parsing the sentiment of contemporary literature.

Keywords: text analysis, sentiment features, CNN-BiLSTM-ATT model, topic model

1. Introduction

The most direct and primitive relationship between human beings and literature is the analysis of literary works, and the aesthetic experience of literary works is the most fundamental spiritual phenomenon of human beings. Reading done in an aesthetic attitude is the important and the most basic reading of mankind, in which the literary work of art and its materialization cease to be an instrument for some other purpose and become the main object of the reader's activity, especially of his consciousness and conscious activity [1-3]. This relationship between human being and the literary work confirms the subjective nature and character of the human being, i.e. the fundamental purpose of this analytical activity is to serve the analyst's own subjective desires. The analyst's love for the work

 Corresponding author.

E-mail address: zhangjingwen1985@sina.cn (J. Zhang).

Received 05 February 2024; Revised 12 April 2024; Accepted 05 December 2024; Published Online 15 April 2025.

DOI: [10.61091/jcmcc127a-188](https://doi.org/10.61091/jcmcc127a-188)

© 2025 The Author(s). Published by Combinatorial Press. This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by/4.0/>).

is the basic prerequisite, and it is characterized by being infected by it, excited by it, and thinking about it. Literary analysis shows the most basic and genuine nature in the relationship between human beings and literature [4-6]. Based on literary analysis, human beings have started literary research activities. Up to now, there are many ways for human beings to study literature, which are more or less the same in China and the West, mainly including the study of literary history, the study of the history of literary trends, the study of writer's theories centered on works, and the literary criticism synchronized with the times, and so on. The study of literary history is a way of study after literature has developed and enriched to a considerable extent. The main task of literary history is to try to accurately describe and summarize the clues and outlook of literary development. The basic way is to arrange a sequence of scattered literary phenomena and writers' works according to certain literary concepts through comparative study, and to give them a specific status in literary history. First of all, evaluation is the main nature of literary history [7-9]. The evaluation of a specific writer in the study of literary history is based on the search and generalization of all the important works of the writer. Secondly, comparison is the principle and basic method of literary history research. The evaluation of any writer is determined by comparing him with other writers. Generally speaking, changes and adjustments in literary concepts, principles of comparison and selection will inevitably raise the demand for rewriting literary history. The study of the history of literary trends identifies continuous trends of thought in literary ideas and creative phenomena that do not appear to be intrinsically connected, describes their direction, and gives historical evaluations and theoretical explanations. Comparison is the basic principle and method of the study of the history of literary trends. The study of the history of literary trends is primarily evaluative and secondarily analytical. The study of writers' theories centered on works is a case-by-case criticism [10]. It takes the detailed analysis and criticism of works as the core, but the final conclusion is based on the positioning and evaluation of writers. Comparison is the basic idea and main method of its research.

Literature [11] proposes two effective features to encode speech information and thus fuse it with textual information. Experimental results on five different Chinese sentiment analysis datasets show that the inclusion of speech features significantly and consistently improves the performance of both textual and visual representations, and outperforms the current state-of-the-art Chinese character-level representations. Literature [12] combines the powerful global feature extraction of the Transformer encoder with the natural sequence feature extraction of GRU. Experimental evaluations were conducted on three real Chinese datasets to reduce the difficulty of Chinese sentiment analysis in natural language processing due to its complexity and uncertainty. Literature [13] proposes a technique suitable for extracting textual information from documents with complex layouts (e.g., newspapers and journals). Prospect analysis and text localization methods are combined. Detailed experiments on two public databases show that mixing layout analysis and text localization techniques improves the page. Literature [14] proposes a web-based text sentiment analysis method based on an improved AT-BiGRU model. The bias of text sentiment analysis due to spelling errors can be effectively avoided and the effectiveness of the improved AT-BiGRU model is demonstrated in terms of accuracy, loss rate and iteration time. Literature [15] proposes a semi-supervised approach where a lasso-based integration model constructs a basic sentiment dictionary from a small training dataset with minimal manual effort, and the sentiment values of the words in the dictionary are propagated to words that do not exist in the training dataset. The performance of this approach is comparable to supervised learning models trained on large datasets. Literature [16] proposes StyleBERT, a new framework for sentiment analysis of textual audio, which enhances the sentiment information of unimodal representations by learning different modal styles, so that the model already obtains valid unimodal representations prior to fusion, thus alleviating the reliance on fusion.

In this study, firstly, by analyzing the syntax and semantics of the text, we propose a text emotion feature extraction method that can be used for computer recognition processing, and classify the text emotion from the emotion lexicon resources, machine learning, and deep learning. Then, data cleaning,

serialization and padding processing are performed on the deactivated words and subwords of the text to improve the efficiency of model training. Then, a CNN-BiLSTM-ATT sentiment analysis model was built based on the feature vectors of the acquired text data. Finally, by comparing the effect of each model in the sentiment analysis of long and short stories, we verified that the model in this paper can effectively identify and extract the expressions of sentiments in contemporary literature in terms of recall, accuracy, precision and loss rate.

2. Emotion Classification and Feature Extraction of Contemporary Literary Texts

2.1. Text Emotion Feature Extraction

Text sentiment feature extraction is mainly through the syntactic syntactic analysis and semantic analysis of the text, from which the text features that express the emotional tendency of the text are extracted. On the other hand, text emotion feature representation is based on the extraction of text emotion features, and transforms the text into feature vectors that can be recognized and processed by computers. Therefore, the related technology of text emotion feature extraction and representation provides a feasible solution for computers to understand and process natural language.

Word2vec mainly contains two kinds of language models, one is the CBOW model which aims at predicting the specified words given the contextual information, and the other one is the Skip-Gram model which predicts the contextual environment of the words given the target words. The CBOW model and the Skip-Gram model can be used in the training process based on the Hierarchical Softmax algorithm or Negative Sampling algorithm. The former is based on Huffman coding tree to maximize the conditional probability from the root node to the leaf nodes: the latter constructs the sample set of positive and negative examples through negative sampling, so as to achieve the improvement of the training speed and the quality of word vectors. In this paper, we take the CBOW model based on Hierarchical Softmax algorithm as an example to briefly explain.

The CBOW model based on Hierarchical Softmax algorithm firstly counts the word frequency of each word in the training corpus, then constructs the Huffman coding tree based on the word frequency, and finally updates the node parameters according to the input data using the gradient ascent algorithm. Figure 1 shows the CBOW model diagram.

The CBOW model contains three layers: input layer, projection layer, and output layer. The inputs to the input layer are the current word w and its context $Context(w)_1, Context(w)_2, Context(w)_3, \dots, Context(w)_{2k}$, where k is the size of the sliding window. The core idea of the CBOW model is to maximize the prediction probability of the current word w , given the context of the current word w . That is, maximize the objective function:

$$L = \sum_{i=1}^{2k} \log p(w | Context(w)_i) \quad (1)$$

The input layer is mapped to the projection layer by doing a cumulative summation of the context $2 * k$ word vectors to get x_w .

$$x_w = \sum_{i=1}^{2k} v(Context(w)_i) \quad (2)$$

$v(Context(w)_i)$ denotes the vector representation of the i nd word.

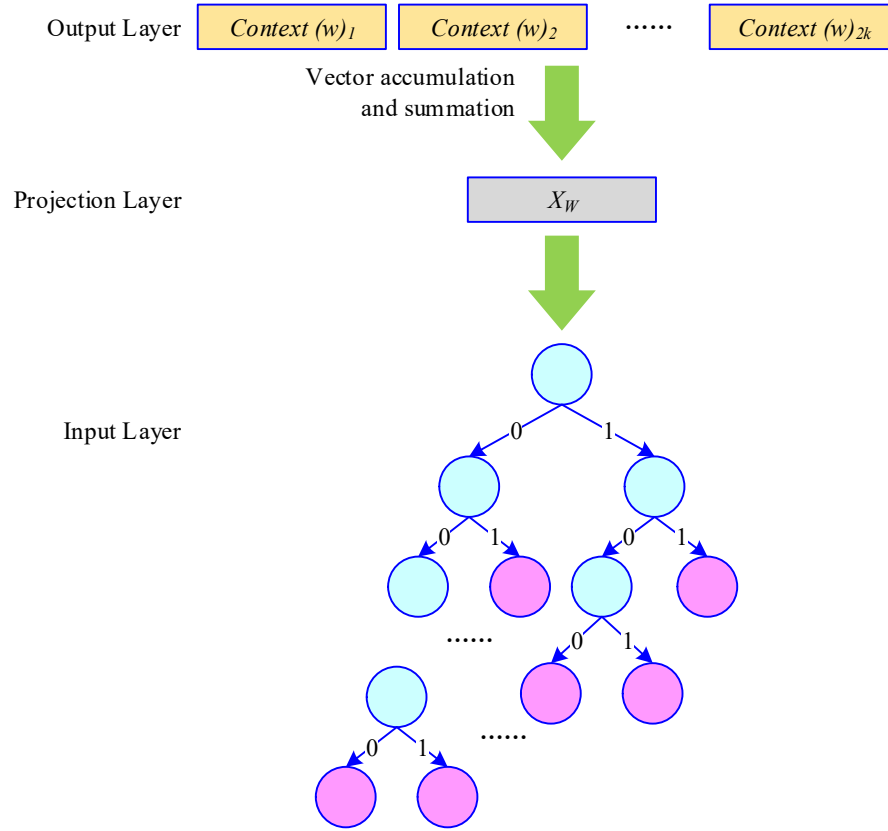


Fig. 1. CBOW model diagram

The output layer constructs a Huffman coding tree based on word frequencies, where the leaf nodes of the tree represent specific words and the inner nodes of the tree consist of Softmax binary classifiers. Therefore, given the current word w , there must exist a one path from the root node to the current word w . Therefore, $p(w|Context(w))$ can be represented as:

$$p(w|Context(w)) = \prod_{j=2}^w p(d_j^w | x_w, \theta_{j-1}^w) \quad (3)$$

$$p(d_j^w | x_w, \theta_{j-1}^w) = \left[\sigma(x_w^T \theta_{j-1}^w) \right]^{1-d_j^w} \cdot \left[1 - \sigma(x_w^T \theta_{j-1}^w) \right]^{d_j^w} \quad (4)$$

l^w denotes the length of the path from the root node to the current word w , d_j^w denotes the categorization label of the j th internal node on the path from the root node to the leaf node, and θ_{j-1}^w denotes the parameters of the $j-1$ th Softmax classifier on the path. So the objective function of CBOW model can again be expressed as:

$$L = \sum_{i=1}^{2k} \sum_{j=2}^{l^w} \left\{ (1-d_j^w) \cdot \log \left[\sigma(x_w^T \theta_{j-1}^w) \right] + d_j^w \cdot \log \left[1 - \sigma(x_w^T \theta_{j-1}^w) \right] \right\} \quad (5)$$

Based on the above objective function, given the training corpus, the target words and the words within the specified window range are used as inputs to the model, and then the gradient ascent algorithm is utilized to update the model parameters so as to obtain the distributed word vector representation of each word.

2.2. Text Sentiment Classification

2.2.1. Text Sentiment Classification Based on Sentiment Dictionary Resources. Sentiment dictionary resources, as an important part of text sentiment calculation, can provide rich feature information for text sentiment analysis. The method of text sentiment analysis based on sentiment dictionary resources mainly extracts the emotion feature words related to emotion expression from the text, and then combines them to calculate the sentiment scores of these emotion words, and judges the emotion tendency of the text according to the scores. The method mainly classifies texts by calculating the text sentiment polarity score.

2.2.2. Machine Learning Based Text Sentiment Classification. The method of text sentiment analysis based on machine learning mainly extracts the sentiment features of the text from the manually labeled training corpus, and then transforms the text data into numerical data that can be understood by the machine, and finally constructs the text sentiment classification model by using the relevant algorithms in the field of machine learning, such as Support Vector Machines, Plain Bayes, and Maximum First Class Models, etc., to realize the sentiment classification of the text. The method first utilizes the idea of plain Bayes to calculate the logarithmic count probability r of each word, as shown in Eq:

$$r = \log \left(\frac{p / \|p\|_1}{q / \|q\|_1} \right) \quad (6)$$

$$p = \alpha + \sum_{i: y^{(i)}=1} f^{(i)} \quad (7)$$

$$q = \alpha + \sum_{i: y^{(i)}=-1} f^{(i)} \quad (8)$$

where $y^{(i)}$ is the label of training instance i , $y^{(i)} \in \{-1, 1\}$. $f^{(i)}$ is the feature vector of training instance i , $f^{(i)} \in R^{|V|}$, V is the feature set, and α is the smoothing factor.

From the formula of r , it can be seen that the value of r represents the sentiment propensity value of each word obtained from learning based on the training corpus. Therefore, r can be used to map the text data into feature vectors and build a binary classification model with support vector machine.

Along with the rise of social network platforms, the study of sentiment analysis of online text data has become a research hotspot. Based on the way of using traditional text modeling, researchers combine the characteristics of network text and introduce new features and methods, so that the existing text sentiment analysis techniques can be applied to network text data.

2.2.3. Deep Learning Based Text Sentiment Classification. Deep learning is a popular research direction in the field of machine learning emerging in recent years, which tries to simulate the learning process of the human brain by building a deep artificial neural network, mainly using a large number of implicit nodes to simulate the neurons in the human brain, so as to achieve the purpose of interpreting and understanding data. At present, the major breakthroughs achieved by deep learning in the fields of image recognition and speech recognition show its powerful learning ability, so more and more researchers apply deep learning to natural language processing and achieve good results.

3. Text analysis system construction

3.1. Text pre-processing

Text preprocessing in natural language processing is a crucial step. To understand the meaning of the text further, we must process the text with word division. Because computers cannot understand and process natural language as well as humans, the text can be converted into mathematical form after word separation, which is convenient for computers to process. At the same time, we are able to improve the training and prediction ability of the model by deleting the deactivated words. Finally, the text is serialized and populated so that it can be used by the model, and the deactivated words and the lexical processing will be elaborated below.

3.1.1. Discontinued words. When preprocessing text, stop words are often removed. Discontinued words, i.e., words that do not contain specific meanings, are common in sentences and usually serve to connect sentence components or to supplement them, and removing them will not have an effect on the expression of sentiment in a sentence. Deactivated words may not accurately reflect a particular mood or meaning, so it is particularly important to remove these deactivated words when conducting sentiment analysis. On the other hand, some keywords with emotional overtones may lead to slower data processing or even greater errors. Removing the deactivated words can also reduce the spatial dimension of the data features and improve the efficiency of model training, etc.

3.1.2. Segmentation. In English texts, words are usually strictly separated from each other. In this way, each word is able to convey the complete meaning. Therefore, when performing word segmentation, we only need to divide the words according to the space, and do not need to use other word segmentation algorithms.

In Chinese text, these words are not clearly distinguished from each other, so we must use specialized algorithms to differentiate them. At present, common Chinese word separation algorithms include: string matching method, probability statistics method, semantic understanding method, the following details of these three methods.

1) String Matching Based Segmentation Method: Words in a large language text base can be accurately categorized by three different ways, i.e., forward matching, inverse matching, and bi-directional maximal matching. First, a complete linguistic text base is needed in order to determine the meaning of each word, then, based on the query results, each unit is accurately classified in order to reach efficient word separation.

2) Segmentation method based on statistical probability: the central idea of this method is to regard each Chinese character as the smallest word unit, which is likely to constitute a phrase if neighboring characters appear frequently in different texts. Therefore, the likelihood of a phrase can be judged by the co-occurrence frequency of neighboring characters. When the co-occurrence frequency is higher than a certain set threshold, it can be assumed that these words form a phrase. Commonly used statistical probability models include N-tuple grammar model, Hidden Markov Model and Maximum Entropy Model.

3) Segmentation methods based on semantic understanding: the main purpose of this technique is to use syntactic and semantic analysis for automatic reasoning in the segmentation process, and to effectively deal with ambiguous situations by combining syntactic and semantic related information. Currently, Python has many available Chinese participle tools, especially worth mentioning is the Jieba participle tool, which is both fast and precise, and easy to use. It provides three main word splitting strategies: exact mode, full mode and search engine mode. Among them, the precise mode can better capture the details of the text, according to the text structure of the best way to split words, provide more accurate results, so in the text analysis and other fields have a wide range of applications. The full mode can usually only divide all the words in the whole sentence and cannot solve the problem of

disambiguation. In contrast, the search engine mode is more powerful, it is an upgraded version of the precise mode, not only can accurately divide the words, but also can be more detailed for the longer words, so as to improve the recall of the words.

3.2. Sentiment analysis model based on CNN-BiLSTM-ATT

The CNN-BiLSTM-ATT model based on BERT is the core of sentiment analysis in this study. The structure of this model is shown in Fig. 2 below, firstly, the input text data is vectorized and represented through Bert's word embedding layer, and then the text data after word embedding is input into the CNN and BiLSTM structure to extract feature vectors respectively. Finally, the feature vectors output from both sides are fused, and the new feature vectors obtained enter the Attention layer to further extract useful feature information using the attention mechanism, and finally enter the Softmax layer for classification.

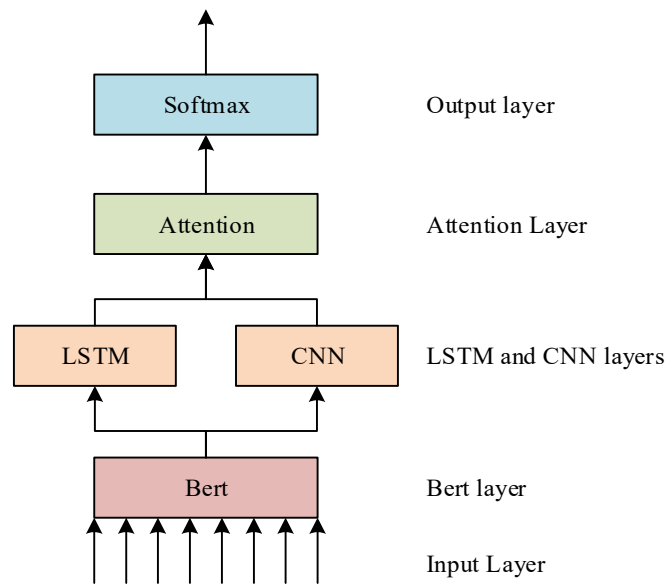


Fig. 2. BERT-CNN-LSTM-ATT model diagram

3.2.1. Input layer. The input layer focuses on text pre-processing of the input text data so that it is easier for the computer to recognize and process this data. First of all, data cleansing is performed on these data, the data obtained from datasets are often dirty data, which means that they are organized directly from the web, and therefore there are certain meaningless words and symbols. The data in it is cleaned by setting some rules through python programming. Then, the data is divided into words, the commonly used word division tool is Jieba word division, which is the most common Chinese word division tool for Python, it combines the dictionary and statistical word division method, the word division efficiency is very high, and you can choose the precise mode to further improve the word division effect, very suitable for text analysis.

3.2.2. BERT layer. The efficiency of BERT for pre-training word vectors is well documented. Here BERT is used as an embedding layer to vectorize the input text, which can greatly understand the meaning of the text before and after the text, making it better to enter the subsequent LSTM and CNN to extract features.

When BERT is used as word embedding, it needs to use a pre-trained model, here this experiment uses Bert_base_chinese model, which is a more suitable pre-trained model for Chinese model.

When using BERT as word embedding method, after reading the data, BertTokenizer, a Bert lexer, is used to lexicalize the text data, which is then converted into Tensor format to facilitate the subsequent data processing work.

3.2.3. LSTM and CNN layers. The vectorized text data embedded from the Bert layer words were fed in parallel to the LSTM layer or CNN layer to obtain the feature vectors of the two models for feature extraction and feature fusion.

Long Short-Term Neural Network LSTM and Text Convolutional Network TextCNN can both extract better text features when processing text data, however, their respective models always lose some information when extracting features, Xingyu Zhao's research found that the current sentiment analysis models are better at combining contextual semantic features and localized features, and therefore proposed fusion of the LSTM and CNN models that Contextual semantic features are extracted using the LSTM layer because LSTM networks are better at processing temporal data, while the CNN model is used to extract local semantic features of the text. Therefore, these two networks are feature fused in this model to further improve their respective text feature extraction capabilities.

3.2.4. Attention layer. The attention mechanism was initially inspired from human visual perception, people will always read things according to the specific things assigned, ignoring unimportant information, so as to maximize the understanding of the information, and greatly improve the efficiency of people's information processing. This unique ability to extract information has been used in deep learning, and has had a wonderful performance in various fields of image processing, speech recognition and natural language processing. The attention mechanism allows deep learning algorithms to focus on localized information, thus improving the recognition ability of the model. The attention mechanism in natural language processing is similar. Attention mechanisms are implemented in natural language processing by allocating different attention to the information acquisition of the whole sentence to further explore the key information and ignore some unimportant content. Initially, the attention mechanism was used in machine translation, but later it has been increasingly generalized to more tasks. The current attention mechanism has become an important research in natural language processing. Its generalized form is a sequence of key-value pairs. It is formulated as follows

$$f(Q, K) = QK^T \quad (9)$$

$$a_i = \text{soft max}(f(Q, K)) \quad (10)$$

$$\text{Attention}(Q, K, V) = \sum a_i V \quad (11)$$

where Q represents a query and $K-V$ is for a key-value pair.

3.2.5. Output layer. The results obtained from the LSTM model and CNN model are weighted with the attention mechanism, and the results are fed into the Softmax layer of this layer for the final classification, which is a four classification task according to the requirements of this experiment. CrossEntropyLoss is added to the model as a loss function to calculate the loss, i.e., the cross-first loss function. This function is very widely used in the classification task. Its formula is shown below.

3.3. Thematic analysis

Topics can be viewed as probability distributions of related words in the study of natural language processing. Topic modeling LDA is an unsupervised machine learning algorithm based on Bayesian ideas. It uses a probabilistic generative model to model the implicit topics of a document, and its basic idea is to de-assume a number of independent implicit topics in the content of the corpus, and based on the probability distributions of these topics all the words of each document of the corpus can be

generated. Thus, documents are understood as distributions of specific implicit topics. LDA is essentially an unsupervised machine learning model that represents the high-dimensional textual word space in terms of the low-dimensional topic word space.

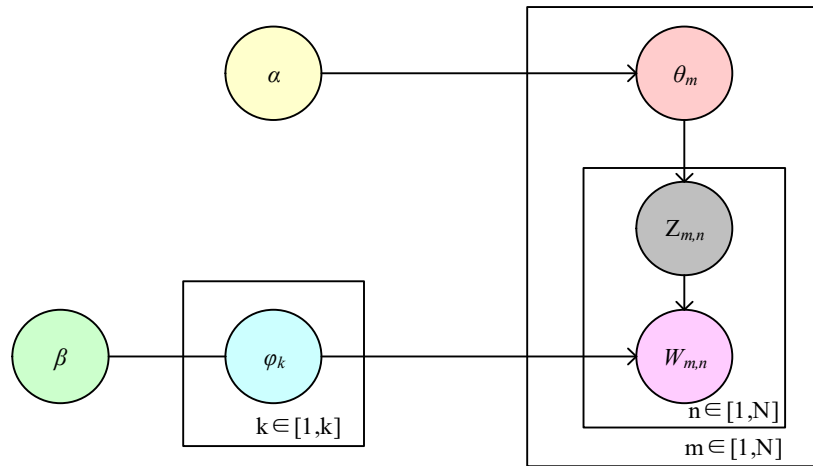


Fig. 3. Basic model of LDA

The basic model structure of LDA is shown in Fig. 3, where M represents the number of documents, K represents the number of topics, and there are V words in the vocabulary, N_m means the number of words in the m th document, $W_{m,n}$ and $Z_{m,n}$ represent the n th word and its topic in the m th document, φ_k represents the probability distribution of all the words in the topic k , and θ_m represents the probability distribution of all the topics in the m th document. θ_m and φ_k obey Dirichlet prior distributions with hyperparameters α and β , respectively.

4. Experimentation and Analysis of Emotional Expression under Textual Analysis

The CNN-BiLSTM-ATT model is compared with other models to verify that the fusion model proposed in this paper has significant advantages and is more suitable for sentiment analysis. This experiment is designed to check the accuracy and effectiveness of this model in terms of various metrics such as accuracy, precision, recall and F1. The corresponding comparison experiments are conducted on seven models such as TF-IDF, Doc2Vec, LSTM (One-hot), BERT-LSTM, BERT-GRU, RoBERTa-LSTM and the algorithm proposed in this paper. Because the contemporary literature text is special, this paper USES the corpus method to build the literary text data set.

4.1. Analysis of long stories

This section is oriented to the long novels in contemporary literature to develop the experimental analysis of emotional expression. The classic long novels in contemporary literature including The Ordinary World, Camel Xiangzi, White Deer Plain and Dust Settles are selected. Based on this, a long story dataset is constructed, which is used to test the analysis effect of different models. The comparison of the experimental results of various models is shown in Table 1. Through the experiments on Data1 and Data2 datasets, it can be seen that the performance of different models on text sentiment analysis, in the NewsData1 dataset, the lowest recall and precision rate is Doc2Vec model, 65.18% and 64.83%, respectively, the model proposed in this paper recall and precision rate, 83.91% and 81.86%, respectively, are higher than the the other six models. In the NewsData2 dataset, the Doc2Vec model has a recall rate of 66.65% and a precision rate of 65.78%, which are lower than

the other six models. The recall and precision of the model proposed in this paper are 84.91% and 82.67%, respectively. It shows that the model in this paper outperforms all other models in sentiment analysis of text.

Table 1. Comparison of model experiment results (unit:%)

Experimental data set	Data1			Data2		
Evaluation index	Recall	Precision	Macro-F1	Recall	Precision	Macro-F1
TF-IDF	72.65	65.81	68.26	73.64	66.48	69.17
Doc2Vec	65.18	64.83	65.14	66.65	65.78	66.02
LSTM(One-hot)	69.42	65.75	66.25	70.26	64.31	67.45
BERT-LSTM	81.71	77.82	79.34	82.01	78.36	80.64
BERT-GRU	83.26	76.91	80.10	83.76	77.41	80.69
RoBERTa-LSTM	82.36	77.64	79.97	82.65	77.68	80.74
CNN-BiLSTM-ATT	83.91	81.86	82.67	84.91	82.67	83.65

The comparison of Macro-F1 values of different models is shown in Figure 4. By plotting the bar chart, the gap between the Macro-F1 values of each model can be clearly demonstrated. In data1, the Macro-F1, of this paper's model is 82.67%, and the Macro-F1 value in data2 is 83.65%, both of which are higher than the results of the other models' sentiment analysis. It proves that the improved CNN-BiLSTM-ATT model fused in this paper has excellent performance on the sentiment analysis task.

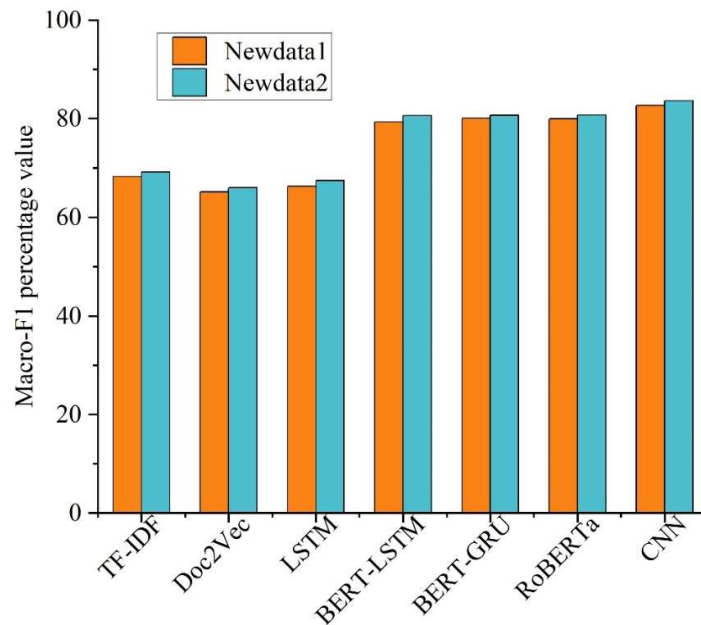


Fig. 4. Different models Macro-F1 value comparison

4.2. Analyzing and Analyzing Short Stories

In this experiment, a dataset constructed from typical short stories in contemporary literature was used for verification. Through text processing of literary short stories such as "The Siege of the City", "Broken Soul Gun" and "The Chess King", a dataset of about 5,612,300 words was built. The dataset contains four kinds of emotion labels, and the sentiment polarity of comments includes positive emotions "gratification" and "joy" and negative emotions "anger", "disgust" and "depression". Each text of a short story is less than 140 kanji in length, which is standard short text data.

The experimental results of each model are shown in Table 2. It can be found through the experiment that in the short text dataset, the model proposed in this paper has the highest precision rate of 98.45%,

and the TF-IDF model has the lowest precision rate of 89.31%. It shows that different models present different analysis effects in different datasets, but the effect of the model proposed in this paper is more stable, with accuracy, precision and recall not less than 98%, which is still better than the other six models.

Table 2. Model experiment results

model	Accuracy rate	loss	Precision rate	Recall rate	F1
TF-IDF	86.51%	41.52%	89.31%	87.44%	81.52%
Doc2Vec	93.43%	41.25%	91.64%	93.16%	91.77%
LSTM(One-hot)	97.43%	40.13%	91.42%	94.31%	93.84%
BERT-LSTM	94.21%	45.76%	91.91%	96.11%	93.42%
BERT-GRU	90.11%	39.74%	90.43%	92.50%	90.63%
4RoBERTa-LSTM	93.64%	36.12%	97.11%	98.15%	95.92%
CNN-BiLSTM-ATT	99.81%	34.11%	98.45%	98.64%	98.57%

The accuracy and loss rate of each model on the short story dataset is shown in Figure 5. The traditional TF-IDF neural network model is at the lowest level of accuracy for sentiment analysis at 86.51%. The model proposed in this paper has the highest accuracy of 99.81%. The BERT-LSTM model has the highest loss rate for text analysis at 45.76% and the model proposed in this paper has the lowest loss rate of 34.11%, which are lower than the other six models. By comparing with the other six models, the accuracy and loss rate of each type of model on the same short story dataset are not the same, and this paper's model has a certain significance of progress, and it also proves that this paper's model can be applied to the text analyzing system, and it has an incomparable superiority compared with the other models.

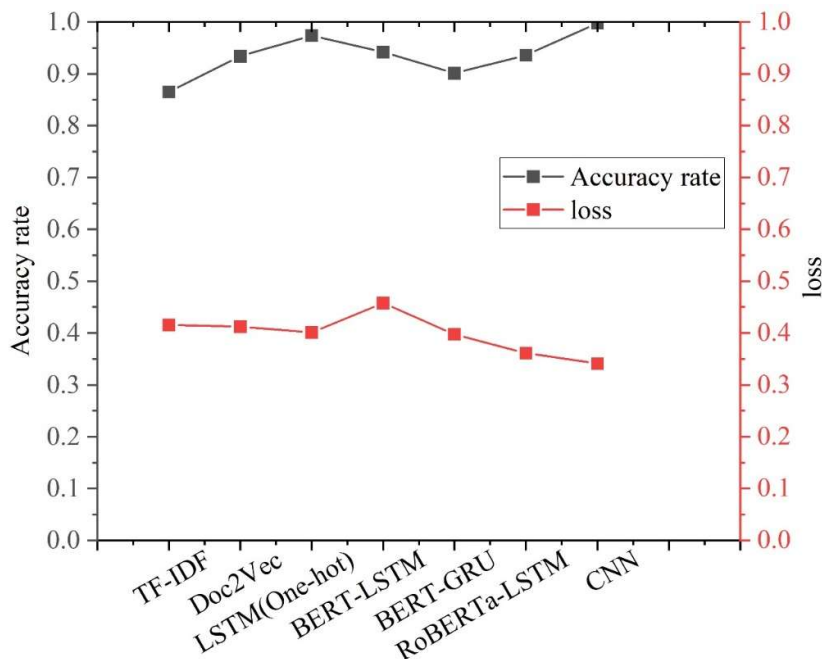


Fig. 5. Accuracy and loss rate of online reviews for each model

5. Conclusion

In this study, a CNN-BiLSTM-ATT sentiment analysis model is built based on the text sentiment feature extraction method, and comparative experiments and analyses are carried out for the emotional

expression of long and short stories in contemporary literature, and the following conclusions are drawn:

- 1) In the Data2 dataset, the recall rate of Doc2Vec model is 66.65% and the precision rate is 82.67%, which are lower than the other six models. The recall rate and precision rate of the model proposed in this paper are 84.91% and 82.67%, respectively. It shows that the performance of this paper's model in the sentiment analysis of the text are better than the other models.
- 2) In the short story dataset, the precision rate of the model proposed in this paper is the highest 98.45%, and the precision rate of the TF-IDF model is the lowest 89.31%. It shows that different models present different analysis effects in different datasets, but the model proposed in this paper is more stable, with accuracy and precision rate not less than 98%, which is still better than the other six models.

In this paper, although the expression of the emotional expression of the literary literature is made by using the text analysis method, the research on the model is not sufficient. Subsequent research can compare the performance of the model in different data sets, and provide the idea for the promotion of the model.

References

- [1] Sumiyoshi, C., Narita, Z., Inagawa, T., Yamada, Y., & Sumiyoshi, T. (2020). M66. facilitative effects of transcranial direct current stimulation on semantic memory examined by text-mining analysis in patients with schizophrenia. *Schizophrenia Bulletin*, 46(Supplement_1), S160-S160.
- [2] K., V., & Gupta, D. (2018). Unmasking text plagiarism using syntactic-semantic based natural language processing techniques: comparisons, analysis and challenges. *INFORMATION PROCESSING AND MANAGEMENT*, 54(3), 394-407.
- [3] Yakut Kili?, I. D., & Pan, S. (2017). Automated domain bias correction and its application in text-based personality analysis*. *International Journal on Artificial Intelligence Tools*, 26(05), 1760023.
- [4] Saura, J. R., Palos-Sanchez, P., & Grilo, A. (2019). Detecting indicators for startup business success: sentiment analysis using text data mining. *Sustainability*, 11.
- [5] Rocha, M. A. D., Morais, P. S. G., Barros, D. M. S., Santos, J. P. Q. D., Dias-Trindade, S., & Valentim, R. (2022). A text as unique as a fingerprint: text analysis and authorship recognition in a virtual learning environment of the unified health system in brazil. *Expert Syst. Appl.*, 203, 117280.
- [6] Petr, H., & Michal, M. (2023). Speech emotion recognition and text sentiment analysis for financial distress prediction. *Neural computing & applications*(29), 35.
- [7] Sun, H., Wang, G., & Xia, S. (2021). Text tendency analysis based on multi-granularity emotional chunks and integrated learning. *Neural computing & applications*(14), 33.
- [8] Corvo, E., & Caro, W. D. (2019). Social network and grief: a text mining analysis. *The European Journal of Public Health*, 29(Supplement_4).
- [9] Caro, W. D. (2020). Infodemia and covid-19: a text mining analysis. *The European Journal of Public Health*, 30(Supplement_5).
- [10] Chen, K. C. (2020). Verbal aggression detection on twitter comments: convolutional neural network for short-text sentiment analysis. *Neural computing & applications*, 32(15).
- [11] Peng, Y. C. E. (2021). Phonetic-enriched text representation for chinese sentiment analysis with reinforcement learning. *Information Fusion*, 70(1).
- [12] Zhang, B., & Zhou, W. (2023). Transformer-encoder-gru (t-e-gru) for chinese sentiment analysis on chinese comment text. *Neural processing letters*.

- [13] Vasilopoulos, N. , & Kavallieratou, E. . (2017). Unified layout analysis and text localization framework. *Journal of Electronic Imaging*, 26(1), 013009.
- [14] Lu, X. , & Zhang, H. . (2021). Sentiment analysis method of network text based on improved at-bigru model. *Scientific Programming*, 2021(12), 1-11.
- [15] Lee, G. T. , Kim, C. O. , & Song, M. . (2020). Semisupervised sentiment analysis method for online text reviews. *Journal of Information Science*, 47(3), 016555152091003.
- [16] Lin, F. , Liu, S. , Zhang, C. , Fan, J. , & Wu, Z. . (2023). Stylebert: text-audio sentiment analysis with bi-directional style enhancement. *Information systems*.