

Exploring Emotion Analysis in Modern Chinese Fiction Using AI Technology

Juan Li¹, 

¹ Yinchuan University of Energy, Yinchuan, Ningxia, 750100, China

ABSTRACT

The emotional curve of a story is the core embodiment of the reading value of a novel, and good novels tend to have similar patterns of emotional changes, which are explored in novels by combining artificial intelligence technology. After collecting modern Chinese novel texts, Chinese word segmentation and de-duplication are performed to complete the novel text preprocessing. In view of the limitations of convolutional neural network (CNN) and recurrent neural network (RNN) in text feature extraction, this paper proposes a multi-channel convolutional and bi-directionally gated recurrent unit (BiGRU) deep learning model, Pt-MCBGA, to mine the emotional polarity in the text and analyze the emotional trend of modern Chinese novels. After a series of comparison experiments, it is demonstrated that the model performance achieves a relatively excellent performance, and the recall rate on the two datasets is improved to 83.53% and 83.69%, respectively. According to the Pt-MCBGA model, the sentiment analysis of the modern Chinese novel *The Legend of the Eagle Shooting Heroes* finds that the novel is dominated by positive sentiment, with both positive and negative sentiment values being relatively high, and that the characters are rich in emotions and have great emotional ups and downs.

Keywords: modern Chinese novel, sentiment analysis, Pt-MCBGA model, convolutional neural network

1. Introduction

The so-called Chinese novels refer to those novel texts written in Chinese written form and its grammatical norms, which are also called narrative works [1-2]. However, when this kind of fiction, which seems to be a narrative, is examined in the specific context of Chinese literary history, it is not difficult to find out a basic fact that the narrative of Chinese novels does not fall from the sky, nor does it grow out of the narrative itself, but it is closely related to the context and even the cultural origins of lyrical works, which are the predecessors of narrative works [3-6].

Even though lyricism has long led Chinese literature, especially ancient literature, lyricism is relative,

 Corresponding author.

E-mail address: 18709511266@163.com (J. Li).

Received 30 May 2024; Revised 15 August 2024; Accepted 10 January 2025; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-070](https://doi.org/10.61091/jcmcc127b-070)

© 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/>).

and without the existence of narrative forms, lyricism would lack support [7-8]. If there is lyricism, there must be narrative, each complementing the other, although lyricism breeds narrative, but narrative constitutes the other half of literature after all [9]. Although the lyric has long been singing Chinese literature, but the lyric has not left the narrative expression for a moment [10]. Whether it is "Poetry Classic", "Chushu", or Tang poetry, Song lyrics and Yuanqu and other lyrical style, almost all of them are lyrical lead singing narrative, the object to attract the feelings of thought, feelings for the Zhi, therefore, lyrical inevitably can not be separated from the narrative, or else, the lyric is not lyrical [11-13]. Therefore, the history of ancient Chinese literature has always been lyric for this, narrative as a supplement, lyric can not be separated from the narrative, the narrative is more dependent on the lyric of the left and right [14-15].

At the end of the Qing Dynasty and the beginning of the People's Republic of China, from Liang Qichao and others advocated the use of "new novels" to "enlighten the people" and "new people's morality", the mode of production of Chinese literature has changed [16]. Here the "new novel" really became "a hybrid form that had to be reinvented at every stage of its development" [17]. However, "the transmutation of the novel's narrative style or other expressive techniques caused by changes in the novel's mode of production (e.g., the production system, newspaper serialization, or bestseller critique) has received little scholarly attention" [18-20]. Today, it seems that academics not only pay little attention to the transformation of the novel's narrative mode, but also pay even less attention to the tradition of lyrical aesthetics in the course of the novel's creative practice, as well as to the methods and systems of poetic criticism that have permeated the history of Chinese literary criticism [21-23].

In this paper, a deep learning model Pt-MCBGA with multi-channel convolution and bi-directional gated recurrent unit (BiGRU) is proposed, which introduces an attention mechanism and a pre-trained word vector model can better extract the feature information and mine the emotional polarity in the text of modern Chinese novels. The pre-trained word vector model is used to directly fill in the high-dimensional word vectors to realize the accurate representation of word vectors. Using multi-channel 1D convolutional neural networks Conv1D and BiGRU to mine text information, more local features and global semantic features of different granularity are extracted, and by introducing the attention mechanism to weight the features, the model is able to quickly obtain more key feature information and improve the possibility of outputting correct probability distributions. Finally, taking modern Chinese novels as an example, the emotional tendency of characters and chapters are analyzed and discussed after text preprocessing and performance testing of the model.

2. Emotional Classification of Modern Chinese Novel Texts

2.1. Preprocessing of Fiction Text

2.1.1. Chinese word splitting. In natural language processing, words are generally regarded as the smallest linguistic unit capable of independent activity. Words are important features for judging the emotional tendency of text, so the accuracy of word segmentation seriously affects the results of text classification. In English, there are natural space characters, which can be relied on to split the text into independent words, while the smallest meaningful granularity of Chinese text is the word, and there is a lack of clear separators between words and words, so it is necessary to do lexical processing before analyzing the Chinese text. Chinese text segmentation refers to the technology of dividing the text sequence into independent and meaningful words.

2.1.2. De-duplication of words. Deactivated words are function words in the text that have no specific meaning, including connectives, adverbs, and personal pronouns, etc. These words generally appear more often in the text and play a minimal role in helping text classification. Deactivating words can effectively reduce the noise of the text and improve the accuracy of text classification. De-duplication removes the words in the deactivation list by scanning the original text. Common Chinese

deactivation word lists include HITC deactivation word list, Baidu deactivation word list, Sichuan University deactivation word list and so on. The actual application can be mixed, but also according to the specific scenario, specifically organize the relevant deactivated words.

2.2. Fiction Text Vectorization

2.2.1. Discrete representation. The models represented by the discrete representation are solo thermal coding, bag-of-words model, TF-IDF model and N-grams model:

1) One-hot coding. One-hot coding is one of the most common methods, and the main steps are divided into two steps: constructing a dictionary after splitting the text into words, and then encoding the words in the text with 0 and 1.

One-hot encoding is easy to understand and simple to construct, but there are some shortcomings, mainly in: the dimension is too high, which is easy to cause a dimensional disaster; sparse matrix, the vector constructed by one-hot encoding has only one dimension with a numerical value; one-hot can't reflect the inter-relationships between words.

2) Bag-of-words model. Bag-of-words model assumes that the words in the text are independent, the bag-of-words model overcomes the dimensionality catastrophe caused by one-hot coding, but it also has the problems of matrix sparsity and the inability to retain the positional information in the sentence.

3) TF-IDF. TF-IDF technique is widely used in the field of data mining and retrieval, TF denotes the word frequency and IDF denotes the inverse text frequency index.

The TF-IDF formula is as follows:

$$TF-IDF(t, d) = TF(t, d) \times IDF(t) \quad (1)$$

where $TF(t, d)$ denotes the frequency of occurrence of word t in text d and $IDF(t)$ is the inverse text frequency index with the following formula:

$$IDF(t) = \log \frac{\text{Total number of articles}}{\text{Total number of articles containing the word } t + 1} \quad (2)$$

Similarly, the TF-IDF model still fails to capture the interrelationships between words.

4) N-grams model. The idea of N-grams model is to judge the next most probable word given a sequence of strings, and in order to preserve the word order in the text, a sliding window operation is added, and N is the size of the sliding window.

The N-grams model is based on the assumption that the occurrence of the n th word in a sentence is only related to the first $n-1$ words and not to the others. The formula is as follows:

$$P(S) = P(w_1, w_2, w_3, \dots, w_n) \quad (3)$$

$$p(w_1, w_2, w_3, \dots, w_n) = p(w_1) * p(w_2 | w_1) * p(w_3 | w_1, w_2) \dots p(w_n | w_1, \dots, w_{n-1}) \quad (4)$$

Introducing the Markov assumption, Eq. (4) can be written as:

$$p(w_1, w_2, w_3, \dots, w_n) = p(w_i | w_{i-m+1}, \dots, w_{i-1}) \quad (5)$$

Among other things, for unigram, there is:

$$P(w_1, w_2, w_3, \dots, w_m) = \prod_{i=1}^m P(w_i) \quad (6)$$

To bigram, there:

$$P(w_1, w_2, w_3, \dots, w_m) = \prod_{i=1}^m P(w_i | w_{i-1}) \quad (7)$$

To trigram, there:

$$P(w_1, w_2, w_3, \dots, w_m) = \prod_{i=1}^m P(w_i | w_{i-2} w_{i-1}) \quad (8)$$

For general natural language processing problems, discrete representation of text is a better choice, but for business scenarios with high accuracy requirements, it is generally necessary to use distributed representation methods for text vectorization.

2.2.2. Distributed Representation. The representation of a word using nearby words in natural language processing is known as distributed representation of text, and distributed representation has become the most commonly used method in modern text representation. Distributed representation of text maps words to vectors of real numbers, a method known as word embedding techniques, also known as word vectors. At present, the most mature text distributed representation models used in academia and industry are Word2Vec model and GloVe model:

1) Word2Vec model. The Word2Vec model has lower algorithmic complexity and can efficiently obtain word vectors. The Word2Vec model contains two basic network structures, namely the continuous bag of words (CBOW) model and the Skip-gram model.

The basic idea of CBOW model structure is to predict the probability of occurrence of a target word w_i given its context as input. As shown in Eq. (9), where S_j is the j th sentence in text T , w_{ij} denotes the i th word in the sentence, and its context is denoted as C_{ij} , θ is the default parameter, and the weight matrix obtained at the end of training is the word vector matrix of the text.

$$\arg \max_{\theta} \prod_j^T \left[\prod_{ij=1}^{S_j} p(w_{ij} | C_{ij}; \theta) \right] \quad (9)$$

The Skip-gram model can be viewed as the inverse structure of the CBOW model, and the optimization objective of the Skip-gram model is to maximize the posterior probability of the text:

$$\arg \max_{\theta} \prod_{w_{ij}}^T \left[\prod_{c \in C_{ij}} p(c | w_{ij}; \theta) \right] \quad (10)$$

2) GloVe model. Using the GloVe model to train word vectors first requires the construction of the word co-occurrence matrix and then the construction of the cost loss function:

$$J = \sum_{i,j}^N f(X_{i,j}) (v_i^T v_j + b_i + b_j - \log(X_{i,j}))^2 \quad (11)$$

where v_i and v_j denote the word vectors of the words, b_i and b_j are the biases of the model, and N is the size of the vocabulary.

2.3. Emotion Classification Process of Modern Chinese Novels

Text Sentiment Classification, also known as Sentiment Propensity Analysis, is one of the most important and most researched topics in the field of Sentiment Analysis [24]. The main process of text sentiment classification includes dividing the training set and test set, and then pre-processing the text, specifically including word division and de-duplication, etc., and then extracting the text features and then carrying out text vectorization, and finally carrying out model training through classification algorithms, and the commonly used classification algorithms include simple Bayesian algorithms and support vector machine algorithms in machine learning, and convolutional neural networks and recurrent neural networks in deep learning. The commonly used classification algorithms include the

simple Bayesian algorithm and support vector machine algorithm in machine learning and the convolutional neural network and recurrent neural network in deep learning.

3. Sentiment analysis model of modern Chinese novels based on AI technology

3.1. Convolutional Neural Networks

Convolutional neural network (CNN) has an expanding range of application areas, from the beginning of the image research to the current speech recognition, natural language processing, etc., with good results, mainly consists of embedding layer, convolutional layer, pooling layer and fully connected layer [25]. Different from the traditional CNN for processing images, its convolution kernel width is determined by the word vector dimension, and it only does convolution in one direction of the text sequence, and sets different sizes of convolution kernels in the convolution layer to extract features, and the formula is shown in Equation (12).

$$FM = f(W \cdot x + b) \quad (12)$$

W in the above equation represents the convolution kernel, x denotes the vector of inputs within the window of the convolution kernel, b is the bias term, and f is the activation function.

3.2. Pt-MCBGA modeling

The Pt-MCBGA sentiment analysis model proposed in this section is shown in Fig. 1. It first undergoes text vectorization representation to convert each Chinese comment text into a standard feature vector matrix, which enables the computer to recognize it successfully: then the local features of the text are extracted using a multi-channel convld and bigru network to capture long-term dependencies between statements; a customized ATTENTION layer is added inside each channel to weight the feature information and obtain a Several independent feature matrices are added, and then they are stitched together to form a large feature matrix; finally, the final probability distribution is output through the fully-connected layer, and the softmax function outputs the final probability distribution after normalizing the vectors.

The pooling layer is introduced for feature filtering and screening, discarding part of the number of non-important features; in order to prevent the interference of noisy features, the Attention mechanism is introduced to weight the feature information, giving it different weights according to the degree of importance, so as to make the probability of the key features to be mined out greater.

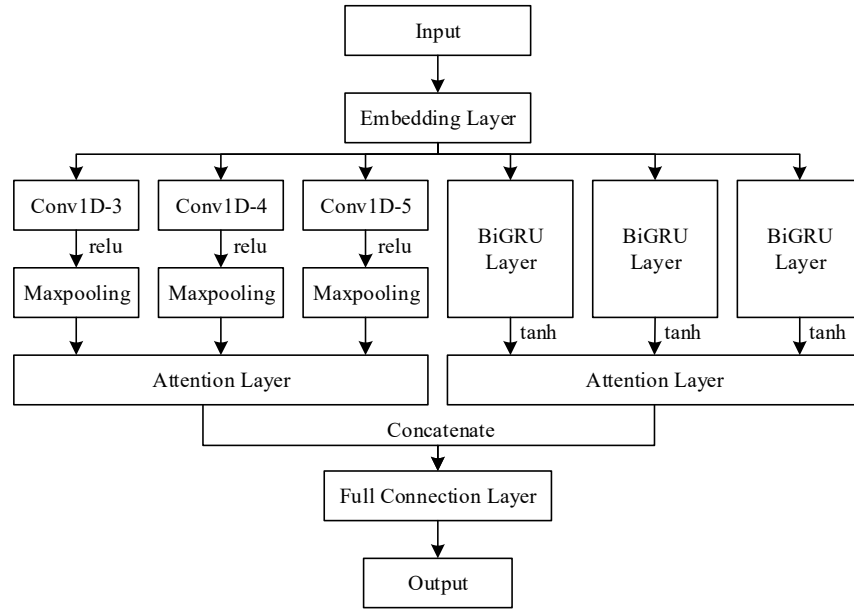


Fig. 1. Pt-MCBGA model structure

3.3. Model structure hierarchy

3.3.1. Word Embedding Layer. Word embedding refers to the mapping of vectors of higher dimensional space with higher dimensionality into lower dimensional space, where each word establishes a connection with a sequence of digits on the real number segment, which is the real sense of what we call word vectors. Since the use of One-hot coding representation produces a large number of sparse matrices with 0 internal elements, and the position of 0 elements represents no feature information, which greatly wastes the memory resource space, we introduced the word embedding matrix. The first 50,000 word vectors from the pre-trained word vector model are populated into the embedding matrix, and then the sparse matrix is multiplied with this matrix to obtain a low-dimensional dense matrix, and the initial vector matrix is downsampled to map the original text content through the Word Embedding layer into a dense matrix composed of low-dimensional vectors.

3.3.2. CNN layer. In this paper, Conv1D is used to extract the internal features of the text, only the width is convolved, and the height is not convolved, and this layer is mainly composed of two major parts, the convolution layer and the pooling layer.

In this paper, we represent the input features as matrix $M \in R^{n \times d}$, where n is the length of the sample word vector, d is the word vector dimension, and assuming that the convolution kernel is $W \in R^{k \times d}$ and k is the width of the convolution kernel, the operation of each convolution is as shown in equation (13):

$$c_i = f(W \cdot M_{i:i+h-1} + b) \quad (13)$$

where $M_{i:i+h-1}$ denotes the sub-matrix of rows i to $i+h-1$ of matrix M , i.e., the vector matrix of words i to $i+h-1$, and the feature mapping vectors obtained by repeated convolution are shown in equation (14):

$$c = [c_1, c_2, \dots, c_{n-k+1}] \quad (14)$$

The features extracted by convolutional computation will be sent to the pooling layer to perform the downsampling operation, and by selecting a fixed-size window for pooling, only the most important feature information in the region will be retained, and the compression and filtration of the vector's

downsampling and features will be accomplished, so that the network parameters can be simplified, the amount of computation can be reduced, and the speed of the model training can be improved, and the phenomenon of overfitting can be avoided. In this paper, we choose the way of maximum pooling, and the related operation is shown in equation (15):

$$c' = \max(c) \quad (15)$$

3.3.3. BiGRU layer. The Gated Recurrent Unit (GRU) introduces the concept of “gates” to control the flow of data, and is equipped with two internal gate control units: the reset gate and the update gate. The reset gate determines how new input information replaces the previous hidden state, and controls the amount of information that is written from the previous state to the current set of candidates; the update gate determines the amount of information that has been brought into the current state from the previous moment, and updates the current hidden information, which controls the extent to which the state information sent from the previous moment is brought into the current state. The update gate determines the amount of information that has been brought into the current state at the previous time, updates the current hidden information, and controls the extent to which state information sent from the previous moment is brought into the current state, and the related internal structure is shown in Fig. 2.

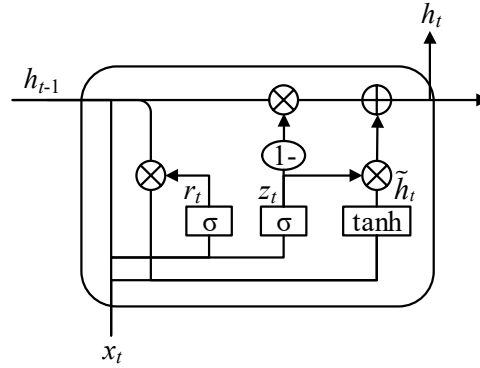


Fig. 2. GRU internal structure

Where x_t is the input at moment t , $\tilde{h}_t$ is the candidate state at moment t , h_t is the hidden state value at moment t , the internal activation function is the tanh function, w_z and w_r correspond to the weight matrices of the update gate and the reset gate, respectively, and b_z and b_r correspond to their biases, respectively, and the related formulas are as follows:

$$z_t = \sigma(\omega_z * [h_{t-1}, x_t] + b_z) \quad (16)$$

$$r_t = \sigma(\omega_r * [h_{t-1}, x_t] + b_r) \quad (17)$$

$$\tilde{h}_t = \tanh(\omega_h * [r_t * h_{t-1}, x_t] + b) \quad (18)$$

$$h_t = (1 - z_t) * h_{t-1} + z_t * \tilde{h}_t \quad (19)$$

The bidirectional gated recurrent unit (BiGRU), on the other hand, is a combination of a pair of forward and backward GRUs, where the two GRU networks extract sequence features in both directions, and then stitch the two output vectors together to get the final result, so that the dependencies between the textual semantics can be fully captured.

3.3.4. Attention layer. *Source* contains numerous $\langle Key, Value \rangle$ data pairs, for a particular feature Q , the correlation between Q and each K is calculated, the weight coefficients of the V

corresponding K are obtained, and then a weighted summation operation is performed on the V to output the final *Attention* score value, which can be abstractly interpreted as equivalent to the importance of this feature. The basic idea can be expressed as the following formula:

$$Attention(Query, Source) = \sum_{i=1}^{L_x} Similarity(Query, Key_i) * Value_i \quad (20)$$

where L_x denotes the length of *Source* and $Similarity(Query, Key_i)$ denotes the dot product computation, i.e:

$$Similarity(Query, Key_i) = Query \cdot Key_i \quad (21)$$

In this paper, a custom Attention layer is added to each channel of multiple CNN and BiGRU networks for feature weighting operation, firstly, the hidden state h_t is nonlinearly transformed using the $\tanh$ function to get a new state λ_t , and then the attention weight α_t is obtained by weighting λ_t , and finally, the hidden state h_t is subjected to a weighted summation operation according to the weight of the attention, to get a new feature vector v_t , which is formulated as follows:

$$\lambda_t = \tanh(w_t * h_t + b_t) \quad (22)$$

$$\alpha_t = \frac{\exp(\lambda_t)}{\sum_{k=1} \exp(\lambda_k)} \quad (23)$$

$$v_t = \sum_t \alpha_t * h_t \quad (24)$$

3.3.5. Full connectivity layer. The main function of the fully connected layer is to map the distributed feature representation obtained from the previous layer to the sample label space, which is simply to act as a “classifier” to integrate the output feature matrix into a single value. The advantage is that it reduces the influence of feature position on the classification result and improves the robustness of the whole network.

In this paper, after the Attention layer, we first use the Concatenate function to splice the feature information extracted from two different structural types of networks into a large feature matrix, and then we connect a fully connected layer afterward, then the standard output vector with a fixed length is continuously passed to the fully connected layer containing two hidden neurons, and then the last layer is densely connected to the two output nodes and normalized by softmax activation function, and then the last layer is connected to the two output nodes. The softmax activation function is normalized and the result is a floating point number between 0 and 1, indicating the distribution probability or confidence, i.e., how likely it is that it belongs to the classification.

4. Emotional Analysis of Modern Chinese Novels

4.1. Word Frequency Analysis of Modern Chinese Novels

4.1.1. Calculation results of feature word weights. By calculating the word frequency and weight of the words after word separation, the top 20 words in terms of word frequency and the top 20 words in terms of weight are selected, as shown in Table 1. Through the analysis of the table, it can be seen that in the evaluation process of modern novels, the words that appear most frequently in the comments are “theme”, “character”, “character”, “lines”, “content”, “story”, etc., indicating that readers pay more attention to the above aspects when evaluating modern novels. Although the word “contagious” is not in the top 10 words of word frequency, its weight has reached 18.49. It shows that the appeal of the novel is the main emotional embodiment, which plays a key role in the analysis of the emotion in the novel.

Table 1. Words in the top 20

Serial number	Vocabulary	Weight	Word frequency	Serial number	Vocabulary	Weight	Word frequency
1	Theme	22.31	267	11	Reality	5.61	142
2	Figures	18.57	255	12	Landscape	6.11	137
3	Character	17.14	229	13	Meaning	4.49	128
4	Lines	27.56	201	14	Art	5.57	122
5	Content	17.08	197	15	Understand	12.39	117
6	Story	5.41	191	16	Reasonableness	9.51	113
7	Branch line	7.28	188	17	Infectivity	18.49	110
8	Vocabulary	8.56	179	18	Uniqueness	3.28	108
9	Image	12.64	165	19	Completeness	6.27	105
10	Plot	20.15	159	20	Effect	2.42	104

4.1.2. High-frequency triads. Taking the modern Chinese novel "The Legend of the Condor Heroes" as an example, Figure 3 shows 20 triples. The 7 triples with the highest frequency (frequency > 50). They are "I won't", "I'm not sure", "she can't", "it already is", "it will be", "except", "will be". The 2 triples that appear more than 40 times are "very" and "she has". Other triples, such as "it will be", "once", "he has been", "impossible", "in the world", appear between 30-40. These triples either represent assertions or negatives. On the whole, these high-frequency triples are dominated by subject-verb structure, prepositional phrases, and conjunctions, and are more colloquial, and less reflect the substantive content of the novel.

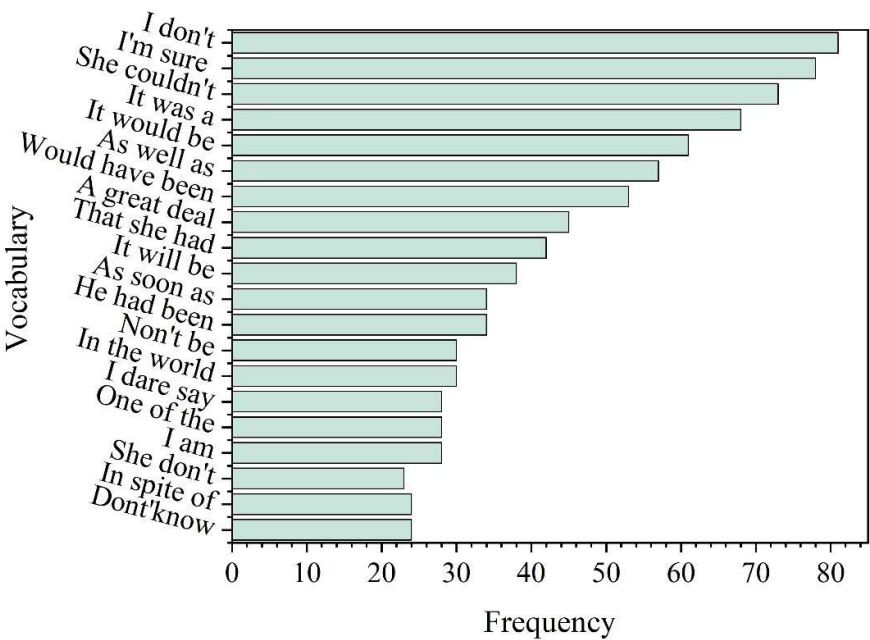


Fig. 3. Maximum frequency triplet

4.2. Text Sentiment Analysis Performance Feedback

4.2.1. Experimental design. The Pt-MCBGA model is compared with other models to verify that the Pt-MCBGA model proposed in this paper has significant advantages and is more suitable for text sentiment analysis. The comparison of experimental models is considered from two aspects: word vector conversion methods and classification models. After careful selection, common models such as

linear regression (LR), LSTM, GRU are chosen as classification models, and then combined with common word vector conversion methods, this paper is designed to carry out corresponding comparison experiments on seven models.

Some parameter configurations of the models during the experiment are shown in Table 2.

Table 2. Part of experimental parameters settings

Name	Numerical value
Cut text	3
Maximum sentence length	516
Hidden layer	516
Discarding rate	0.3
Learning rate	1e-4
Batch size	51
Training set dictionary size	31201

4.2.2. Analysis of experimental results. The experimental results are shown in Table 3. Through the experiments on NewsData1 and NewsData2 datasets, the performance of different models on sentiment analysis can be seen, proving that the Pt-MCBGA model constructed in this paper performs well on text sentiment analysis, and the model has a recall of 0.8353, a precision of 0.8150, and an average F1 value of 0.8250 on NewsData1, which is better than the other models on NewsData2 with a recall of 0.8369, precision of 0.8254 and average F1 value of 0.8311, outperforming the other models.

Table 3. Comparison of model experiment results (unit: %)

Experimental data set	NewsData1			Newsdata2		
Evaluation index	Recall	Precision	Macro-F1	Recall	Precision	Macro-F1
TF-IDF	72.4	65.1	68.55	73	66	69.32
Doc2Vec	65.39	64.29	64.84	66.1	65.44	65.77
LSTM(One-hot)	69.19	62.63	65.75	69.82	63.88	66.72
BERT-LSTM	81.28	76.85	79	81.6	77.7	79.6
BERT-GRU	82.99	76.54	79.63	83.38	76.95	80.04
RoBERTa-LSTM	81.8	77.26	79.46	82.1	77.41	79.69
Pt-MCBGA	83.53	81.5	82.5	83.69	82.54	83.11

4.3. Analysis of Emotional Tendencies in Modern Chinese Novels

4.3.1. Character Frequency View. The Character Frequency View represents the frequency of occurrence of each character role in each chapter of a novel in the form of a matrix heat map, using colors to express the information. In this paper, we use this visual representation to fill the grid using gradient colors to give the information among the colors, and indicate the weights or frequencies through the change of color shades, which is concise and intuitive. There are usually many different color schemes used to illustrate heat maps, and in practice gradients of color can enhance the representation of details.

In this paper, the gradient of blue is used, and the frequency of each character in each chapter is indicated by color according to the statistical results, in the heat map matrix, each grid on the horizontal axis indicates the content of one chapter in the novel, and the vertical axis is divided according to the extracted characters, with different frequencies encoded as different color depths, as shown in Fig. 4.

The representative chapter text content of Legend of the Eagle Eagles and the extracted 28 main

characters are selected, and the frequency of each character is obtained by calculating the word frequency, and the heat map drawn based on the character frequency can facilitate the user to quickly locate the source of the character's appearances and find out the aggregation density of the characters in the content of each chapter. In order to avoid visual confusion, the frequency values of the characters' appearances are filtered by setting a threshold value to avoid visual confusion caused by the representation of characters with too low relevance.

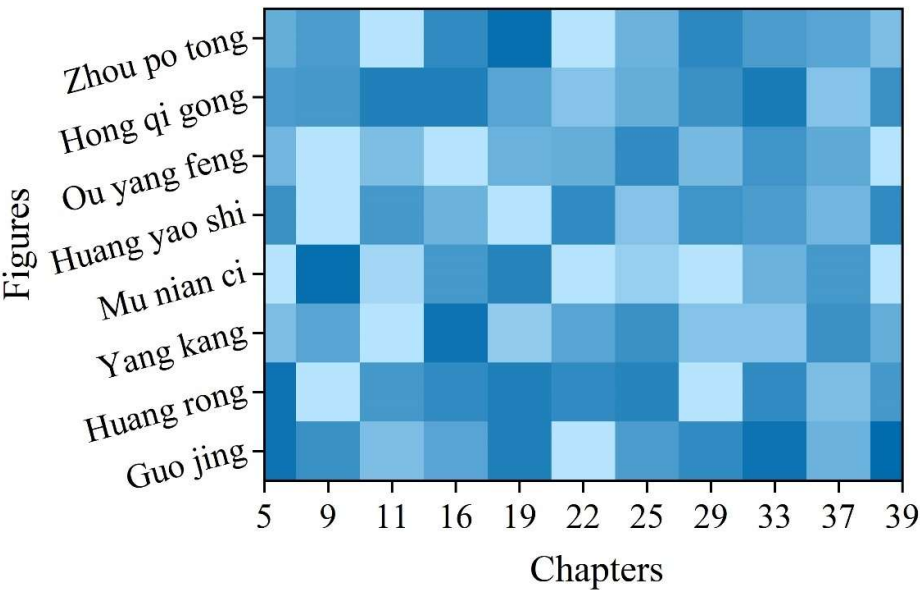


Fig. 4. Character frequency view

4.3.2. Character Emotion Analysis. Sentiment analysis using the Pt-MCBGA model, due to exploratory experiments, only the sentiment lexicon and the sentiment positive and negative corpus that comes with the toolkit were used in the model training, so there are still errors in sentiment assignment, and some of the sentence sentiment calculations are not in line with the cognition, which may in turn have an impact on the overall sentiment change.

Based on the sentiment analysis of the sentiment dictionary, the sentiment word list with detailed scores produced by Tsinghua Yuanbo is selected as the base word list, and the method of assigning scores to positive and negative sentiment words and adjusting the weight and polarity of negative words is adopted to perform the calculation of the sentiment value.

Figure 5 shows the emotion change graph of Guo Jing, the main character of “The Legend of the Eagle-Shooting Heroes”. From the graph, there are big ups and downs in the character's emotional changes, and in the 22nd chapter his negative factor emotion increases dramatically, and the emotional score is low. In the 34th chapter, the main character's inner emotion is high, and the emotion score is relatively high.

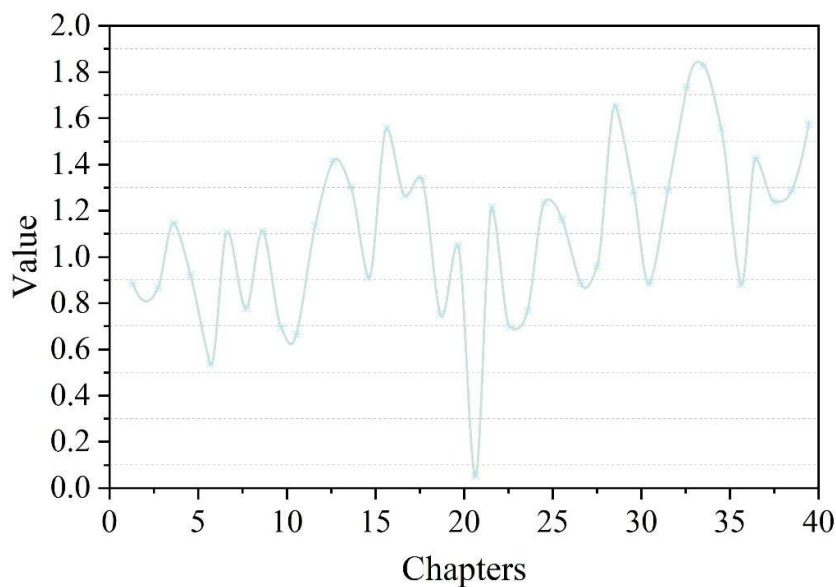


Fig. 5. Emotional curves based on emotional dictionaries

4.3.3. Chapter Sentiment Analysis. There are 40 chapters in the novel “The Legend of the Eagle Shooting Heroes”, and the emotional changes in each chapter are different, with positive values indicating positive emotions and negative values indicating negative emotions, and the larger the absolute value, the more distinctive the emotional differences are, and the emotional changes in each chapter are shown in Figure 6.

From the figure, the positive emotion values in the novel are much more than the negative emotion values. Chapter 29 of the novel has the largest emotion difference value, more than 120, and this result is closely related to the ups and downs of the plot in the novel. Chapter 33 is the lowest emotional low in the novel. Overall, the graph of the change in emotion of each chapter in the figure shows the difference in the emotion value of each chapter. However, where the difference in emotion values is relatively small does not actually mean that emotions do not fluctuate or are absent.

Figure 7 shows a graph of positive versus negative emotions for each chapter. The graph shows that the highest point of both positive and negative affective values are in the last chapter, chapter 40, indicating that the affective change of the text reaches its highest point in the last chapter. The positive affective values of all the chapters are above 100, and the positive affective values of chapters 6, 18, and 40 are even above 200. Meanwhile, the negative affective values of all the chapters are greater than or equal to 60, except for chapters 24 and 36, and the negative affective value of chapter 40 reaches above 180.

The great changes in positive and negative affective values reflect the great emotional ups and downs of the novel's characters and emphasize the theme of the novel.

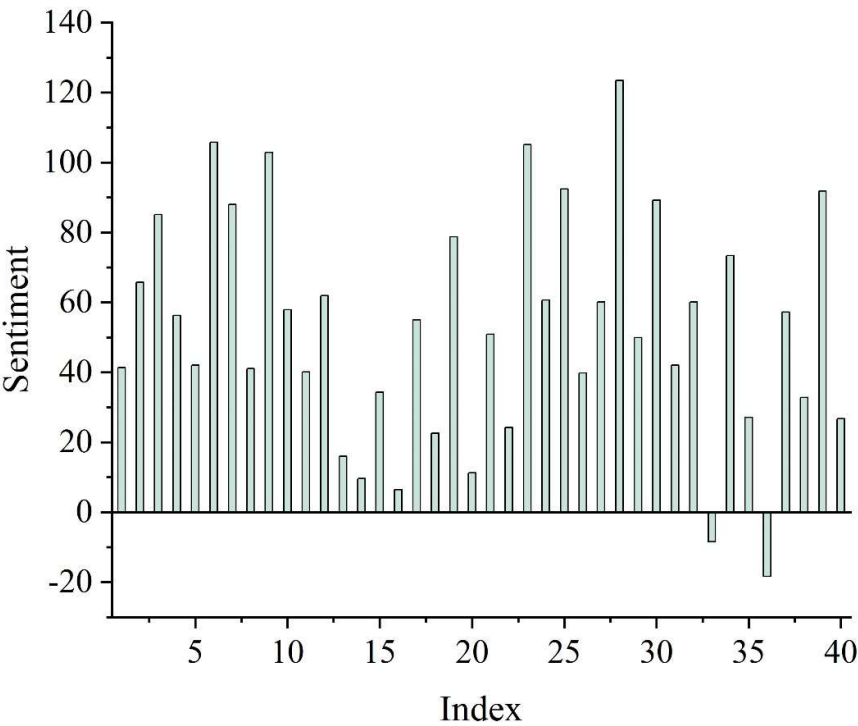


Fig. 6. Different types of emotional changes

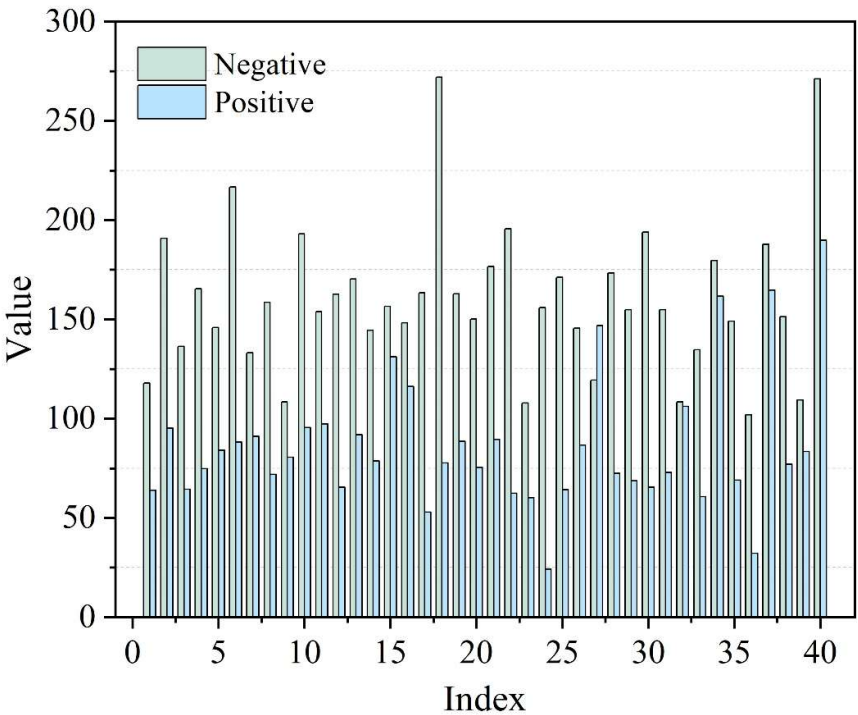


Fig. 7. Positive and negative emotions

5. Conclusion

In this paper, a Pt-MCBGA deep learning model is proposed, and 7 groups of fusion models are used to conduct comparison experiments with the model to determine that the Pt-MCBGA deep learning model is more suitable for text sentiment analysis, and it is proved through the comparison experiments that the recall rate of the model on NewsData1, NewsData2 is 0.8353, 0.8369; precision

rate is 0.8150, and 0.8254; the average F1 value is 0.8250, 0.8311, the performance effect is better than other models.

Sentiment analysis of modern Chinese novels based on Pt-MCBGA model, including character sentiment analysis and chapter sentiment analysis. Taking Legend of the Eagle-Shooting Heroes as an example, the novel is dominated by positive emotions, with both positive and negative emotion values being relatively high, and the characters are rich in emotions and have great emotional ups and downs. It fully reflects the ups and downs of the characters' emotions in modern Chinese novels, and also makes the plot of the novel more exciting.

References

- [1] Wolf, W. (2020). Lyric Poetry and Narrativity: A Critical Evaluation, and the Need for "Lyrology". *Narrative*, 28(2), 143-173.
- [2] Kjerkegaard, S. (2021). A Lyrical 'I' Beyond Fiction: Yahya Hassan and Autobiographical Poetry in Denmark After Karl Ove Knausgård's *My Struggle*. *The European Journal of Life Writing*, 10, 75-95.
- [3] Hillebrandt, C., Klimek, S., Müller, R., Waters, W., & Zymner, R. (2017). Theories of Lyric. *Journal of Literary Theory*, 11(1), 1-11.
- [4] Rodriguez, A. (2017). Lyric Reading and Empathy. *Journal of Literary Theory*, 11(1), 108-117.
- [5] Maslov, B. (2018). Lyric universality. *The Cambridge Companion to World Literature*, 133-48.
- [6] Culler, J. (2017). Extending the Theory of the Lyric. *Diacritics*, 45(4), 6-14.
- [7] Caserio, R. L. (2017). Novel Narratives, Lyric Flights. *Diacritics*, 45(4), 44-66.
- [8] Jancsó, D. (2019). *Twentieth-century metapoetry and the lyric tradition*. De Gruyter.
- [9] Fitzgerald, W. (2023). *Catullan provocations: lyric poetry and the drama of position* (Vol. 1). Univ of California Press.
- [10] Müller, R., & Stahl, H. (2021). Contemporary Lyric Poetry in Transitions between Genres and Media. *Internationale Zeitschrift für Kulturkomparatistik*, 2, 5-24.
- [11] Damayanti, P. A. D., Soethama, P. L., & Udayana, I. N. (2023). An analysis of Taylor Swift's song lyric the man using feminist literary criticism theory. *Langua: Journal of Linguistics, Literature, and Language Education*, 6(1), 81-88.
- [12] Váña, J. (2020). Theorizing the social through literary fiction: For a new sociology of literature. *Cultural Sociology*, 14(2), 180-200.
- [13] James, D. (2017). In Defense of Lyrical Realism. *Diacritics*, 45(4), 68-91.
- [14] Barziev, O. (2021). LYRICAL INTERPRETATION OF "I" IN THE POETRY OF ANBAR OTIN, DILSHODI BARNO AND SAMAR BONU. *CURRENT RESEARCH JOURNAL OF PHILOLOGICAL SCIENCES*, 2(06), 61-66.
- [15] Wu, S. (2020). *Modern archaics: continuity and innovation in the Chinese lyric tradition, 1900-1937* (Vol. 88). BRILL.
- [16] Culler, J. (2017). Lyric Words, not Worlds. *Journal of Literary Theory*, 11(1), 32-39.
- [17] Stahl, H. (2021). The 'Novel in Poems'—An Emerging Genre. *Internationale Zeitschrift für Kulturkomparatistik*, 2, 87-117.
- [18] Hempfer, K. W. (2017). Theory of the lyric: a prototypical approach. *Journal of Literary Theory*, 11(1), 51-62.

- [19] Zettelmann, E. (2017). *Discordia concors. Immersion and artifice in the lyric*. *Journal of Literary Theory*, 11(1), 136-148.
- [20] Culler, J. (2018). *Narrative Theory and the Lyric*. *The Cambridge Companion to Narrative Theory*, 201-16.
- [21] Bekhta, N. (2017). *Emerging Narrative Situations: A Definition of We-Narratives Proper*. *Emerging Vectors of Narratology*, 101-26.
- [22] Tiffany, D. (2020). *Lyric poetry and poetics*. *Oxford Research Encyclopedia of Literature*.
- [23] Garofalo, D. M. (2019). *Victorian Lyric in the Anthropocene*. *Victorian Literature and Culture*, 47(4), 753-783.
- [24] Chao Guoqing, Liu Jingyao, Wang Mingyu & Chu Dianhui. (2023). *Data augmentation for sentiment classification with semantic preservation and diversity*. *Knowledge-Based Systems*.
- [25] Yunji Zhao & Jun Xu. (2024). *A small sample bearing fault diagnosis method based on novel Zernike moment feature attention convolutional neural network*. *Measurement Science and Technology*(6).