

Research on Comprehensive Evaluation Model of Cultural Value of Tang Dynasty Literary Works Based on Convolutional Neural Network and Text Feature Extraction

Jia Liu¹, 

¹ Ministry of Culture and Education, Pingdingshan Polytechnic College, Pingdingshan, Henan, 467000, China

ABSTRACT

In the process of social development of Tang Dynasty, literary works behind the depth of interpretation and expression, systematized spiritual concepts. In this paper, the text data of Tang Dynasty literary works are processed by word division and de-discontinued words, and it is intended to use Transformer model to realize the word vector transformation of text data, and put the word vector into Text-CNN network for iterative training to realize the text feature extraction. By means of text feature screening, the cultural value assessment system of Tang Dynasty literary works is formed, and a comprehensive evaluation model of cultural value is designed under the role of convolutional neural network and text features, and using the model of this paper, the cultural value of Tang Dynasty literary works is assessed. The accuracy rate of cultural value classes “II”, “III” and “V” is 1, while the accuracy rate of cultural value classes “I” and “IV” have accuracy rates of 0.98 and 0.96, indicating that the model in this paper can accurately assess the cultural value in Tang Dynasty literary works.

Keywords: transformer model, convolutional neural network, textual features, evaluation models

1. Introduction

In the context of continuous social and economic development, people are in the fast-paced life, coupled with the penetration of foreign cultures, which makes people's values change a little, some people's efforts to strive for the purpose is no longer a social responsibility, ideals and beliefs, but more for the pursuit of materialism, which leads to their own impatience. At the same time, scientific and technological progress and informationization have brought a huge amount of entertainment and consumer culture, which further affects the inheritance and development of the excellent traditional culture, and many people begin to neglect the pursuit of the spiritual level, at which time they need to enrich the spiritual world [1-3]. Culture and tradition are important factors that determine the values,

 Corresponding author.

E-mail address: 13937511861@163.com (J. Liu).

Received 03 February 2024; Revised 12 April 2024; Accepted 10 November 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-359](https://doi.org/10.61091/jcmcc127b-359)

© 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/>).

behavioral orientation and artistic interests of a society, while ancient literature is an important carrier of traditional culture [4]. Ancient Chinese literature reflects the social environment and spiritual outlook of a specific era, and contains the personal consciousness, values and emotional attitudes of the creators [5]. Through reading and analyzing ancient Chinese literary works, inheriting and innovating them, it is believed to be of great significance in shaping good character and passing on the national spirit, as well as providing valuable cultural and moral resources for contemporary society [6-8]. For example, the vast number of literary works and the completeness of literary styles in the Tang Dynasty are unmatched by previous generations, and their cultural values have been artistically reproduced in a variety of literary styles [9]. In the context of the new era, in order to let more people perceive the charm of ancient Chinese literature, it is necessary to implement the inheritance work in depth, and by introducing the measures of evaluating the cultural value of literary works, it can effectively revolutionize the presentation pattern of ancient Chinese literature [10-11]. At the same time, with the help of modern technological means such as convolutional neural network and textual feature extraction, the Chinese spirit, Chinese wisdom and Chinese values in the literary works are objectively embodied, and at the same time, the multiple values of cognition, society and culture are disseminated [12-13].

This project uses relevant technology to obtain the cultural value information of Tang Dynasty literary works, and develops text pre-processing means such as text segmentation and deletion of stop words, so as to make its text data meet the research standards. At present, the commonly used text representation model consists of bag-of-words model, Word2vec model, and Transformer representation model respectively. Considering the research content and the priority of the model, we intend to use the Transformer representation model to complete the word vector conversion of the text data, so as to facilitate the text feature extraction of Text-CNN model. Based on the requirements of text feature importance standard, complete the construction of cultural value assessment index system and model. Finally, we synthesize the dataset to validate and evaluate the model of this paper.

2. CNN-based model for assessing the cultural value of literary works

2.1. Text pre-processing

Since most of the text data of Tang Dynasty literary works on the Internet are disorganized and contain a lot of noise information, in order to get the text data that meets the task of text categorization, it is necessary to clean these raw text data. So the first step of the text characterization process is to preprocess the text data. The most important part of the text preprocessing stage is text segmentation and de-duplication.

2.1.1. Text Segmentation. The semantic expressions in the text are represented by a combination of words, and these words usually contain rich textual semantic information and Tang Dynasty literary value, so in the text pre-processing stage, the first thing to do is to perform lexical processing on the text data [14]. Currently, the methods of word separation for text data are roughly divided into three categories. The first category is based on dictionary matching, and the text data is divided into words by matching the professional dictionary. The second category is based on probabilistic statistical disambiguation, which disambiguates text through statistical correlation probability. The third category is based on text semantic participle, which analyzes the semantic information of text data and then participle it.

2.1.2. Removal of stop words. Deactivated words in text are words that interfere with the text categorization task, such as conjunctions and intonations that appear in large numbers in the text. The accuracy of text classification can be greatly improved by de-duplication. Many mature deactivation word lists have been developed, and researchers only need to delete the interfering words in the original text data according to these perfect deactivation word lists.

2.2. Text representation

Text data written based on human habitual thinking may become difficult to understand for computers, so it is necessary to adjust the content of the text through the text representation method, so that it is converted into a form that is easy for computers to process. At present, the commonly used text representation models are bag-of-words model, Word2vec model, Transformer representation model, considering the research content of this paper, Transformer representation model is used. In the traditional neural network model (e.g., CNN or RNN) structure, the text sequence needs to be trained in a specific order from front to back or back to front, but such a training mode because the current computation is very dependent on the results of the previous time period of the computation of the computation of the model can not carry out parallel computation, and at the same time, the model is prone to gradient disappearance during the work of the model, so that the semantic information is lost. To address the above problems, researchers propose a completely different model based on the attention mechanism from the traditional network model, the Transformer model, which utilizes the attention mechanism to represent the distance between any element in the text sequence as a constant, which not only solves the problem of gradient disappearance that the traditional neural network model is prone to, but also has good parallelism [15]. The structure of the Transformer model is the classic encoder and decoder structure, and in practice, researchers usually use multiple encoders and decoders in a stack. The first step of the model is to do vector embedding of the text sequence that needs to be input, use the encoder to encode the completion of the encoding, and then decode it with the previous output into the decoder, and finally get the probability of the next word through the Softmax function. Next this paper will introduce the Transformer model each module in detail.

2.2.1. Input layer. The input layer of the Transformer model inputs both the word embeddings of the text sequence and the corresponding positional encoding. The model generally uses pre-trained vectors such as Glove or Word2Vec as word embeddings. Because the Transformer model processes all the words in the text sequence at the same time, it cannot capture the sequence information of the text semantics as RNN, in order to prevent the loss of the position information of the word embeddings in the encoding, the researchers added position encoding to the input of the Transformer model, and there are various methods for calculating position encoding, such as absolute position encoding, relative position encoding, etc. The most appropriate method is chosen according to the needs of different real-world scenarios. The most suitable encoding method is selected according to the needs of different practical scenarios.

2.2.2. Encoders. The encoder in the Transformer model is divided into two modules, a multi-head self-attention module and a feed-forward neural network module, with a residual linkage module and a layer normalization module connected to each of the output positions of these two modules. The multi-head self-attention module, on the other hand, consists of a combination of multiple self-attention layers, which are computed in parallel, allowing each head to get focused on different representational features, thus allowing the model to have a larger capacity. The formula for the attention mechanism is as follows:

$$Atten(Query, key, Value) = Atten(Q, K, V) \quad (1)$$

Assuming that the number of self-attention layers in the multi-attention mechanism is a , the model is projected to W , Q , and V by a linear transformation with the following equation:

$$head = Atten(QW^Q, KW^K, VW^V) \quad (2)$$

$$MultiHead(Q, K, V) = Con(head_1, head_2, \dots, head_a)W^O \quad (3)$$

Where, W^V , W^Q , W^K are the parameters of a head linear transformation in the multi-head attention mechanism, and the parameters spliced after the multi-head output result are W^O .

The Transformer model is computed using the self-attention mechanism, assuming that the input text sequence is x , and the output result of the multi-head attention mechanism is Q' , K' , V' . In order to be able to compute the attention value faster, the attention computer result of the Transformer model is computed by scaling the clicks with the following formula:

$$Atten(Q', K', V') = \text{soft max} \left(\frac{QK^T}{\sqrt{d_k}} \right) V' \quad (4)$$

where d_k is the dimension of the Q -value and the K -value, by dividing by d_k so that the result obtained by dot-multiplication is not too large. The feedforward neural network module consists of a ReLU activation function linear activation function whose main function is to perform nonlinear transformations.

The residual connection and layer normalization module consists of two parts: residual connection as well as layer normalization. Residual linking allows the model to focus on the difference between the two layers by feeding the sequence of unprocessed information from the previous layer directly into the next layer. Layer normalization, on the other hand, normalizes the activation values of the model, increasing the training speed of the model and allowing it to converge faster.

2.2.3. Decoders. The model architecture of the decoder is relatively similar to that of the encoder, with the difference that it uses a masked multi-head attention layer instead of the multi-head attention layer in the encoder. The role of the mask in the model is to mask certain selected values so that the masked values are invalidated when the model updates the parameters. The role of masking in the Transformer model ensures that the model does not go out of range when performing the prediction task.

2.2.4. Output layer. After the work of the model decoder is done, the output vectors of the decoder will be fed into the Linear Transform and Softmax layers. The linear transformation layer maps the vector output from the fully connected neural network model into a Logits vector, and finally the Softmax function converts the Logits vector into the probability of occurrence of the sequence, and selects the word with the highest probability as the output at this point.

2.3. Text-CNN-based text feature extraction

The Transformer model embeds the corresponding word vector data into the word embedding layer of the CNN model after word vectorization of the fetched document. The model mainly consists of four parts: input layer, convolutional layer, pooling layer and fully connected layer, in which the input layer is also called word embedding layer. The input layer of the model needs to input a fixed-length text sequence, and the length of an input sequence is specified by analyzing the length of the samples in the corpus set L . Sample sequences shorter than L need to be filled, and sequences longer than L need to be intercepted. Ultimately, the input layer is fed with word vectors corresponding to the individual words in the text sequence.

The input to the convolutional layer is a matrix representing sentences with dimension $n \times d$, i.e., there are n words per sentence and each word is represented by a d -dimensional word vector. Assuming that $X_{i,j+1}$ represents X_i to X_{i+j} words, using a convolution kernel W with width d and height h to perform a convolution operation with $X_{i,j+h-1}$ (h words), and then activated with an activation function to obtain the corresponding features c_i , the convolution operation can be expressed as:

$$c_i = f(w \cdot x_{i:j+h-1} + b) \quad (5)$$

After the convolution operation, a $n-h+1$ -dimensional vector c shaped like:

$$c = [c_1, c_2, \dots, c_{n-h+1}] \quad (6)$$

Maximum pooling is used in the pooling layer of the Text-CNN model, i.e., it reduces the model parameters and ensures that a fixed-length fully-connected layer is obtained as an input to the output of the non-fixed-length convolutional layer. The core role of the convolutional layer and the pooling layer in the classification model is feature extraction, from the input fixed-length text sequence, the primary features are extracted using the local word order information, and the primary features are combined into the high-level features, and the feature engineering step in traditional machine learning is eliminated by the convolutional and pooling operation. In this paper, we mainly use this model to do dimensionality reduction of features.

2.4. Comprehensive Evaluation Model of Cultural Value of Tang Dynasty Literary Works

2.4.1. Indicators for the evaluation of cultural values. A Text-CNN-based text feature extraction model is used to determine the cultural value evaluation index system of Tang Dynasty literary works, the main principle of which is to rank the importance of the extracted text features, and screen out the text features with feature importance greater than 0.8 as the cultural value evaluation index system of Tang Dynasty literary works, so as to lay a theoretical foundation for the subsequent comprehensive evaluation model.

2.4.2. Standardized processing. In order to reduce the impact of outliers, this paper uses data normalization to process the data, calculated by the formula:

$$X' = \frac{x - \bar{x}}{\sigma} \quad (7)$$

Where x is the original value, $\bar{x}$ is the mean and σ is the standard deviation.

In practice, the Standard Scaler in sklearn.preprocessing library can be called directly for normalization.

2.4.3. Determining the model structure. The structure of the comprehensive evaluation model is shown in Fig. 1, and this study uses a multi-channel input layer containing a fixed amount of indicator features based on the classical Lenet network. The input layer receives one-dimensional evaluation data as input, and the convolution layer extracts local features of the evaluation data by sliding one or more one-dimensional convolution kernels. The pooling layer is used to reduce the dimensionality of the output of the convolutional layer, thus reducing the number of model parameters, and this model uses maximum pooling, where the maximum value within the pooling window is output as the feature value. The fully connected layer is to arrange the feature vectors obtained from the last layer into one-dimensional feature vectors by rows or by columns as the fully connected layer, and the expert's evaluation of the cultural values in the Tang Dynasty literary works, etc. as the output layer.

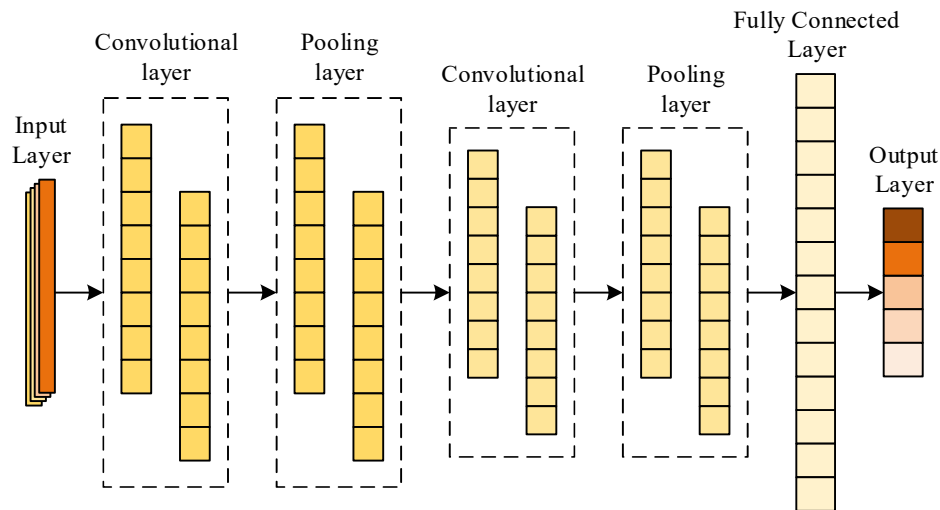


Fig. 1. Comprehensive evaluation model structure

2.4.4. **Model Training and Testing.** The data subset is partitioned into training set, validation set and test set in the ratio of 6:2:2. Among them, the training set is used for training modeling, the validation set is used for tuning the hyperparameters of the model and performing preliminary evaluation, and the test set is used for evaluating the generalization ability of the model. In this paper, we adopt the k -fold cross-test method, where the training set is randomly partitioned into k independent subsets, and only 1 of them is used at a time, while the remaining $(k-1)$ subsets are used as the practice set, and this procedure is repeated until all the subsets have been selected. The final performance results can be obtained by averaging the models from the k training sessions.

3. Exploration of the cultural value of Tang Dynasty literary works

3.1. Text feature extraction analysis

3.1.1. **Description of the data set.** This paper collects 14,884 ancient poems from the literary works of the Tang Dynasty, which originated in 618 and ended in 907, a period during which a large number of excellent literary works emerged, and the statistical analysis of the research dataset is shown in Figure 2. The ancient poems collected in this paper come from the website of Ancient Poetry, although it doesn't cover all the data of ancient poems, but it basically includes all the recorded, more complete, and more studied poems of 20 poets of the Tang Dynasty. As can be seen from the chart, the largest number of literary works in the Tang Dynasty is Bai Juyi (3084), followed by Du Fu (1881), and then Li Bai (1184), with a large gap between the remaining poets, which will not be expressed in detail, but will be based on the data in the chart.

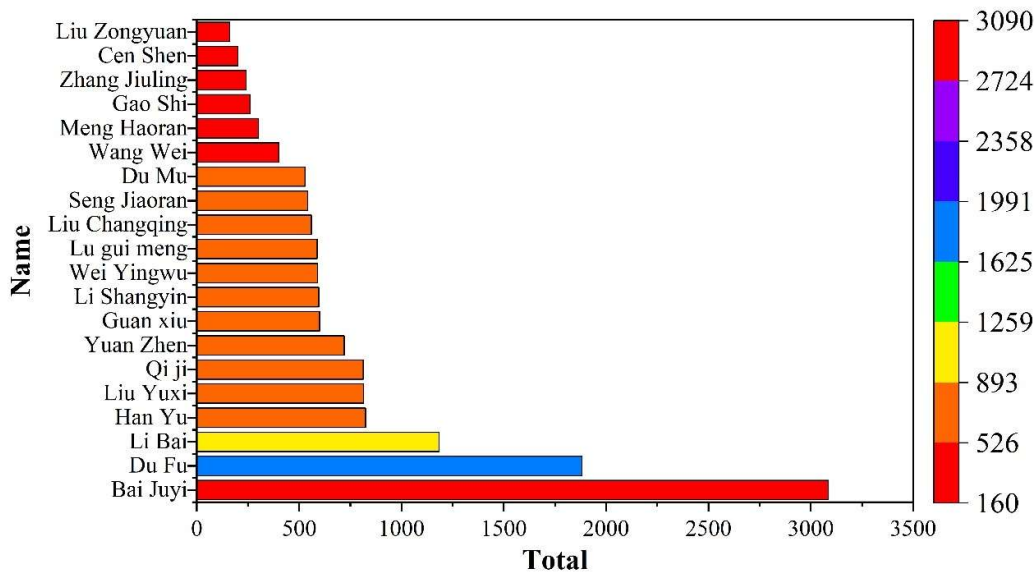


Fig. 2. Research data set statistical analysis

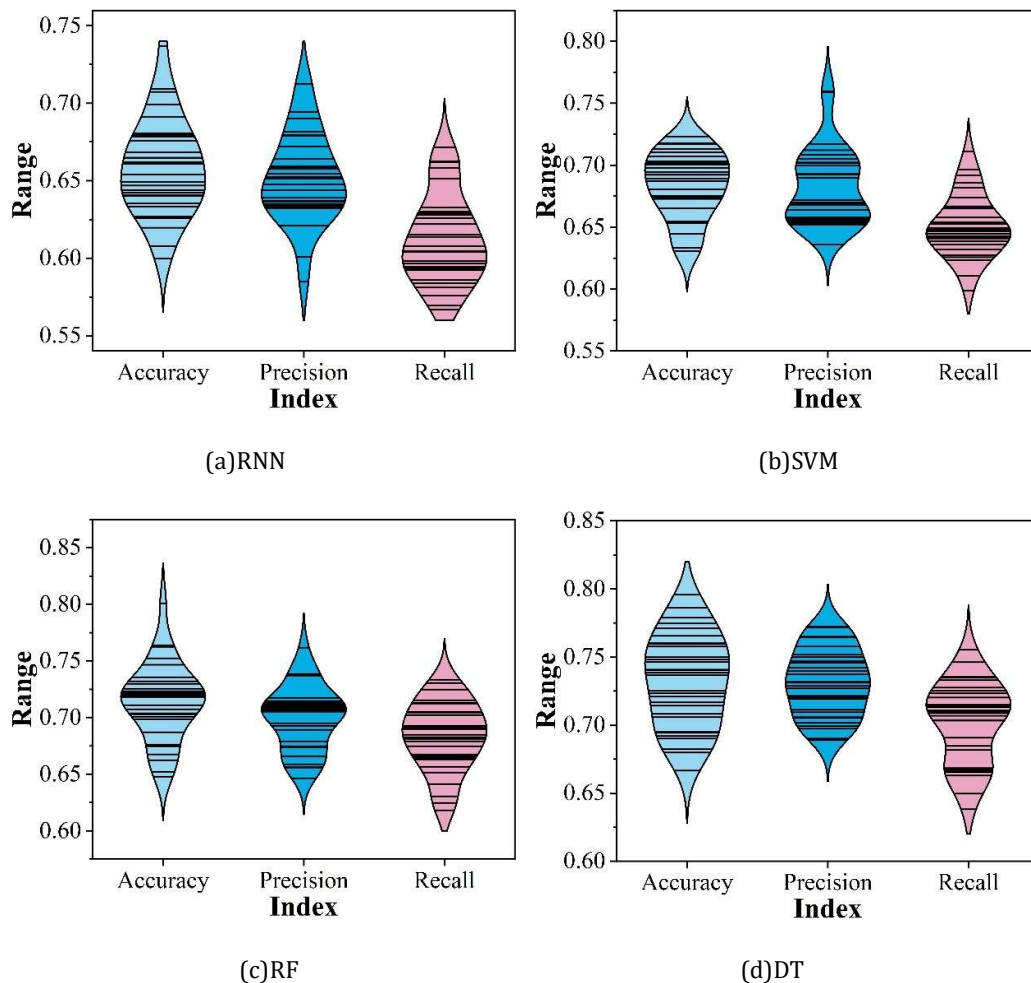
3.1.2. Training effect of TextCNN with different hyperparameters. Then Text CNN is used as the basic deep neural network model, and this paper utilizes tensorflow for the training of Text CNN. Firstly, the relevant hyperparameters of Text CNN are experimentally configured. In the text categorization task, the feature word vectors in the input text feature representation matrix have strong correlation, if the size of the convolution kernel is too small, the extracted word representation information is incomplete, while if the size of the convolution kernel is too large, it is easy to introduce noise. Therefore, with the help of the text classification task convolutional kernel size is generally taken as an empirical value of 5. In this paper, two different sizes of convolutional kernel combinations (3,4,5) and (4,5,6) are set to carry out experiments on the convolutional layer, respectively.

The parameter Batch_size is generally set to be a finger multiple of 2, because whether it is a multi-core CPU or GPU acceleration, the memory management is still based on bytes as the basic unit to do hardware optimization, and the finger multiple of 2 setting will effectively improve the hardware processing efficiency of operations such as matrix slicing and tensor computation. Moreover, the CNN model is very sensitive to the adjustment of the batch size, this paper adopts the manual adjustment method, and according to the training results, this paper will set the iteration number to 10. In order to make the established model reach the optimal as far as possible, the results obtained by the model in the case of using different parameters are compared, and the best parameter in terms of classification results is selected. Table 1 shows the accuracy of the training and prediction sets for some of the different hyperparameter values. In this paper, the convolutional kernel size is set to (4, 5, 6), Batch_size is set to 64, and Num_filter is set to 64, and TextCNN is trained to classify the above research dataset, and the classification results achieved are relatively good overall.

Table 1. Some TextCNN training results under different hyperparameters

Convolution kernel size	Batch_size	Num_filters	Training set accuracy rate	Verification set accuracy	Test set accuracyrate
3,4,5	8	64	0.562	0.628	0.673
3,4,5	16	64	0.674	0.733	0.557
3,4,5	32	64	0.629	0.643	0.694
4,5,6	64	64	0.767	0.739	0.794
4,5,6	128	64	0.532	0.506	0.654
4,5,6	256	64	0.641	0.641	0.588

3.1.3. Evaluation of Text Feature Extraction Effect of Different Models. For the problem of text feature extraction of Tang Dynasty literary works, the performance of the model is usually measured by accuracy, precision and recall, accuracy is used to assess the overall ability of the model recognition, and accuracy, precision and recall are the three commonly used indexes for evaluating the goodness of the text feature model. The evaluation of the text feature extraction effect of different models is shown in Fig. 3, where (a)~(f) are RNN (Recurrent Neural Network), SVM (Support Vector Machine), RF (Random Forest), DT (Decision Tree), DNN (Deep Neural Network), and the method in this paper, respectively. From the data performance in the figure, this paper's method is better than the other five methods in the range of 0.75~1.00 for each index, which verifies that this paper's method has an excellent application effect in the text feature extraction in the Tang Dynasty literary works, and makes the results of the subsequent research more persuasive.



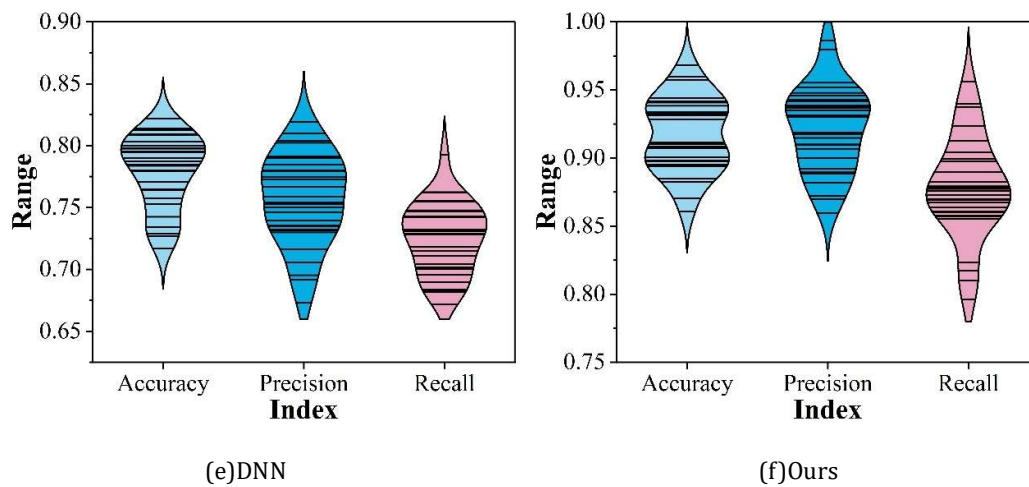


Fig. 3. Evaluation of text feature extraction effect of different models

3.2. Comprehensive evaluation model analysis

3.2.1. Determination of the system of evaluation of cultural values. Through the research and analysis above, it can be seen that the text feature extraction model based on TextCNN has excellent performance. On this basis, the model is used to obtain 15 cultural value indicators of Tang Dynasty literary works, and the construction of the cultural value evaluation index system is completed by calculating its importance value, and the cultural value evaluation system is shown in Table 2. Through the data performance in the table, it can be seen that in the acquisition of 15 indicators, only the humanistic spirit manifestation A5 (0.721), style diversity B5 (0.762), the depth of the personal state of mind exposed C5 (0.703) does not meet the requirements of the indicator importance standard, A5, B5, C5, these three indicators to do the elimination process, and finally determine the cultural value of the Tang Dynasty literary works of the indicator system, the system consists of 12 secondary indicators and 3 primary indicators.

Table 2. Cultural value evaluation system

Primary index	Symbol	Secondary index	Symbol	Importance value	Result
Ideological connotation	A	Depth of philosophical speculation	A1	0.91	Reserve
		Social insight strength	A2	0.816	Reserve
		Historical record significance	A3	0.955	Reserve
		Moral education function	A4	0.849	Reserve
		Display of humanistic spirit	A5	0.721	Delete
Literary aesthetics	B	Charm of language art	B1	0.838	Reserve
		Aesthetic sense of rhythm	B2	0.829	Reserve
		Image-building level	B3	0.958	Reserve
		Height of artistic conception creation	B4	0.983	Reserve
		Style diversity	B5	0.762	Delete
Emotional expression	C	Love theme delicacy	C1	0.988	Reserve
		Friendship shows sincerity	C2	0.986	Reserve
		Nostalgia is contagious	C3	0.885	Reserve
		A strong sense of home and country	C4	0.834	Reserve
		Revealing depth of personal mood	C5	0.703	Delete

3.2.2. Model empirical analysis. Jupyter Notebook is applied to train the established comprehensive evaluation model of cultural values for multiple learning, and the confusion matrix in statistics is used to measure the classification performance of the model. Confusion matrix is a visualization tool, also known as error matrix, is a specific two-dimensional matrix, mainly used to compare the classification

results and the actual measured values, the accuracy of the classification results can be displayed inside a confusion matrix, each column represents the predicted category, the total number of each column indicates the number of data predicted to be in that category, each row represents the true category of the data, the total number of data in each row indicates the category number of data instances, and the values of the elements on the diagonal indicate the number or proportion of correct classifications. The confusion matrix for the model training set is shown in Figure 4 below. According to the numerical values of the main diagonal of the confusion matrix, the classification accuracy of cultural value grades "I.", "III.", "II." and "V." is 1, indicating that all the four cultural value grades are correctly predicted, the classification accuracy of grade "IV." is 0.99, and the model has a probability of 0.01 that it is incorrectly predicted as grade "I."

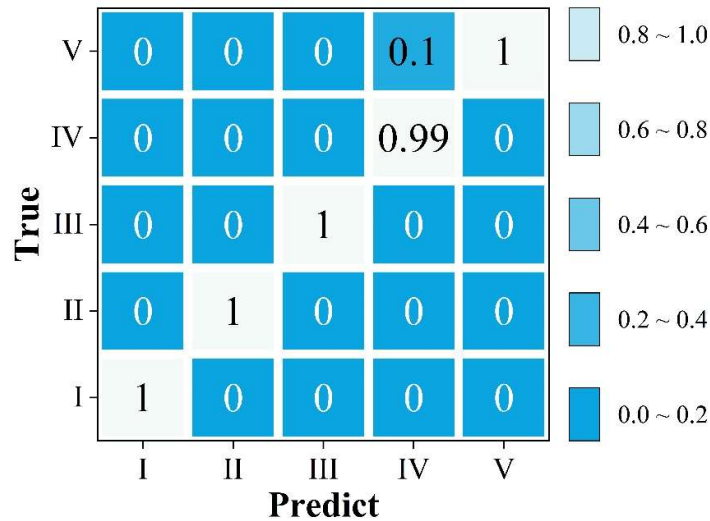
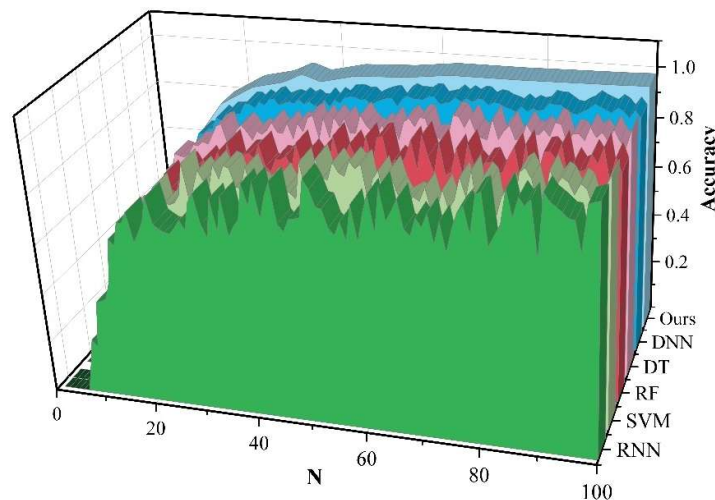


Fig. 4. Confusion matrix of model training set

Figure 5 shows the loss function and accuracy change curve of the model, where (a)~(b) are the accuracy and loss function, respectively. It can be seen that, compared with other models, this paper's model performance is more excellent, with the gradual increase in the number of iterations, this paper's evaluation model in the training set of the loss function continues to decline tends to 0, the number of iterations is 100 when the loss function is 0.0561, the accuracy rate continues to improve gradually tends to 1, in the 100th time of learning the model accuracy rate is 99.65%. Overall, the trained CNN evaluation model has a high accuracy rate, the model learning ability is strong, the model will be converted to json.h5 format to save, easy to follow up on its testing and analysis.



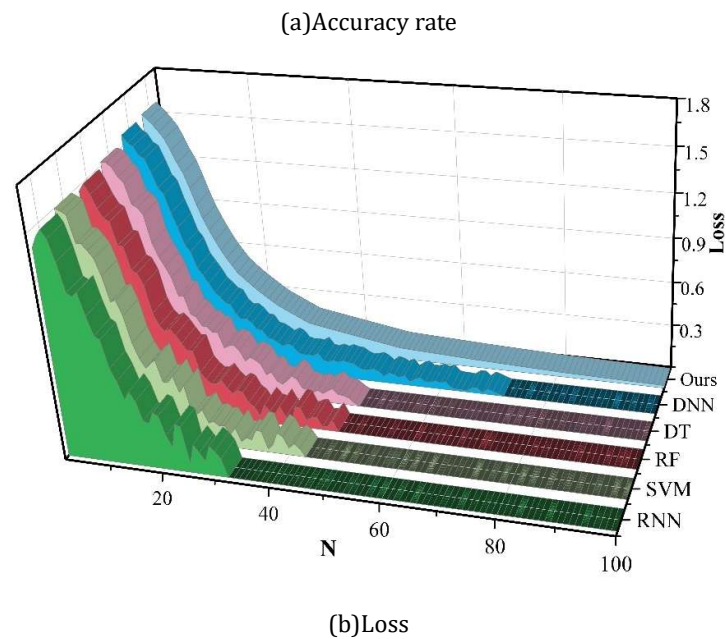


Fig. 5. Training set accuracy and loss function change curve

In terms of personnel factors and operating environment, relevant professionals were invited to score 12 indicators such as depth of philosophical discourse, degree of social insight, and significance of historical records. A total of 450 groups of test samples based on the CNN evaluation model were loaded into the saved convolutional neural network evaluation model, and Table 3 shows the first five groups of data to be evaluated displayed by Jupyter Notebook. After the evaluation results output after the learning of the network model established in this paper, the network evaluation accuracy is 99.56%.

Table 3. The first five set of evaluation sample data

N	A1	A2	A3	A4	A5	B1	B2	B3		C5	Lv.
1	2	2	2	4	1	2	4	3		2	I
2	4	4	4	5	3	5	5	5		5	II
3	4	2	2	3	5	2	4	3		1	III
4	4	2	3	3	4	3	3	4		1	IV
5	4	1	3	4	2	2	1	4		5	V

The obfuscation matrix for the network validation output is shown in Figure 6. The numerical values of the main diagonal of the confusion matrix show that the classification accuracy of cultural value levels "II.", "III." and "V." are all 1, indicating that all the three cultural value levels are correctly evaluated, the classification accuracy of grade "I" is 0.98, the model has a probability of 0.02 to incorrectly evaluate it as grade "I", the classification accuracy of grade "IV." is 0.96, and the model has a probability of 0.04 to incorrectly evaluate it as grade "V.". From the overall situation, it is considered that the evaluation results are in line with the actual situation, and the cultural value evaluation model based on convolutional neural network is scientific and reasonable, and some cultural values of grades "I" and "IV" may have changed due to various factors, so they are developing towards unsatisfactory levels.

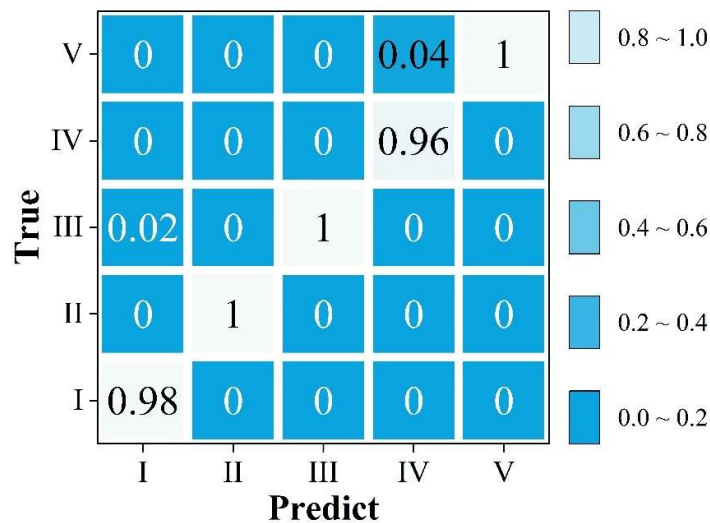
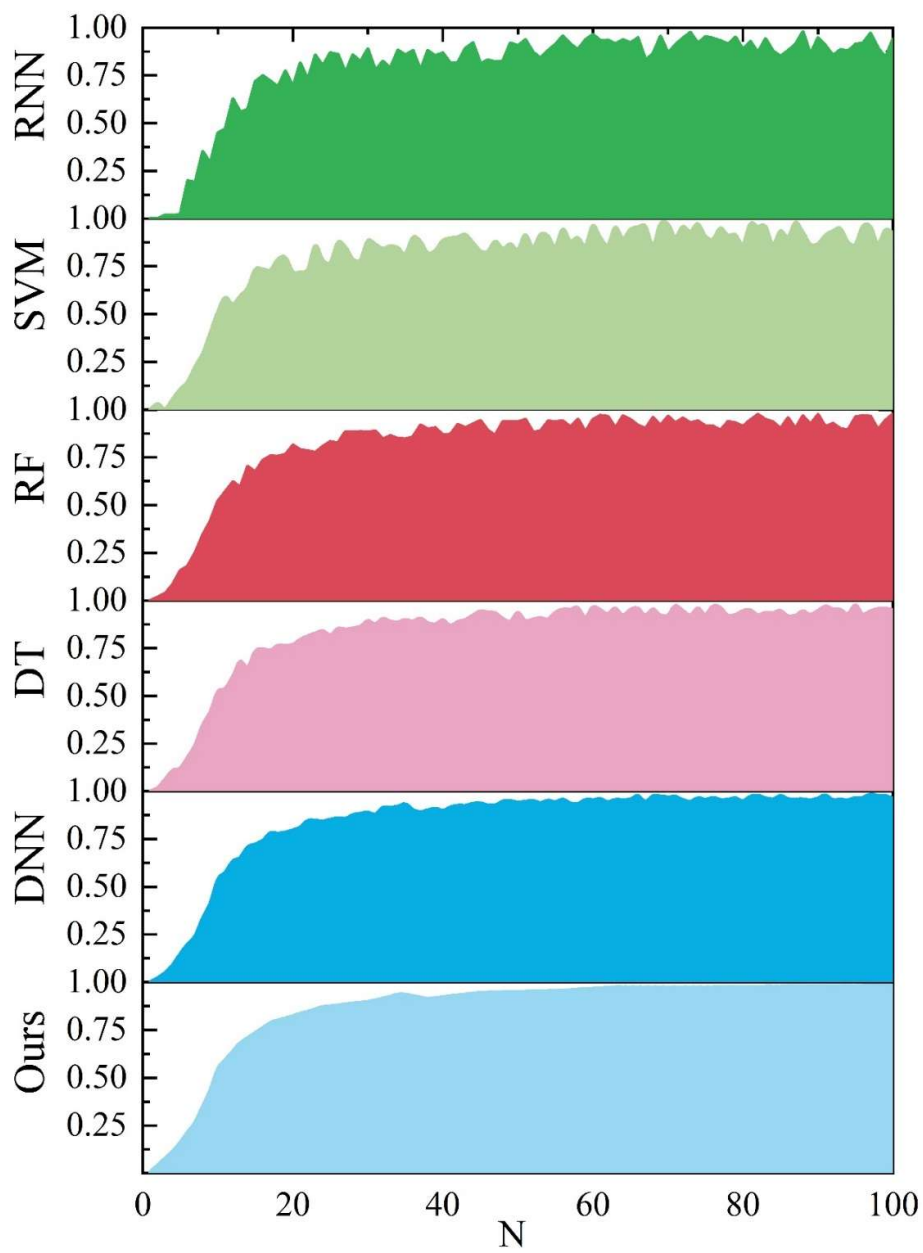
**Fig. 6.** Confusion matrix of CNN model test set

Figure 7 shows the accuracy and loss function change curve of the test set, it can be seen that with the gradual increase of the number of iterations, the loss function of the evaluation model on the test set continues to decline converging to 0, and the accuracy continues to improve gradually converging to 1. Compared with other models, the learning ability and accuracy of the cultural value evaluation model based on CNN and textual feature extraction are higher, and it has a certain value for cultural applications.



(a)Accuracy rate

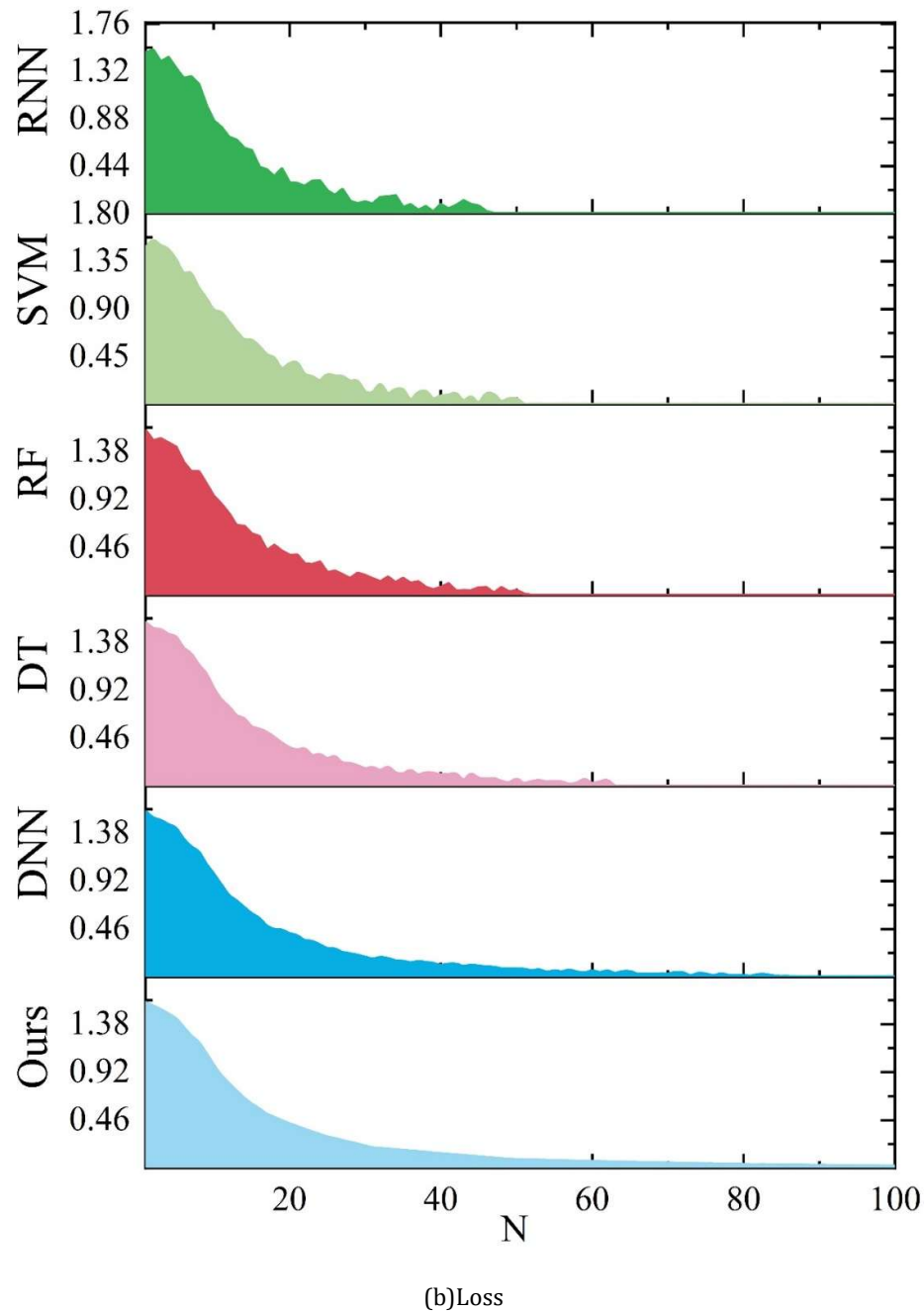


Fig. 7. Test set accuracy and loss function change curve

4. Conclusion

In this paper, for the interference information contained in the text data of Tang Dynasty literary works on the Internet, the text data are preprocessed, and after word vectorization is completed with the help of Transformer model, the corresponding word vector data are embedded into the word embedding layer of CNN model. After that, the Text-CNN model is used to complete the text feature extraction, and the importance degree of the extracted text features is calculated, and the features with importance degree greater than 0.8 constitute the cultural value assessment index system. On this basis, text features and CNN are synthesized to build a comprehensive evaluation model of cultural value, and the model is used to explore and analyze the cultural value of Tang Dynasty literary works. Under the conditions of the five grades set, the classification accuracy of the evaluation model of this paper for the cultural value grades “II”, “III” and “V” is 1, while the classification accuracy

of the other two grades is 0.98, 0.96, 0.98, 0.96, 0.95 and 0.95 respectively. 0.98, 0.96, indicating that the evaluation model constructed in this paper can accurately capture the cultural value of Tang Dynasty literary works, and can carry out reasonable classification and assessment of cultural value, while the overall evaluation results are in line with the actual situation, and have a guiding value for the dissemination and inheritance of traditional culture.

References

- [1] Xiang, Y. (2018). Folk culture: the cornerstone of Chinese culture's influence. *International Communication of Chinese Culture*, 5, 301-304.
- [2] Zhou, Y., & Liu, G. (2024). A Study on the Inheritance and Transformation of Ancient Literary Traditions in Online Literature. *The Educational Review, USA*, 8(10), 1205-1208.
- [3] Zhao, Z. (2023, August). A Study of the Meaning of Ancient Chinese Literature in the New Era. In 4th International Conference on Language, Art and Cultural Exchange (ICLACE 2023) (pp. 299-303). Atlantis Press.
- [4] Cai, Z. Q. (2015). Sound over ideograph: the basis of Chinese poetic art. *Journal of Chinese Literature and Culture*, 2(2), 545-572.
- [5] Yao, S. (2022). AN ANALYSIS OF THE CURRENT SITUATION OF ANCIENT CHINESE LITERATURE FROM THE PERSPECTIVE OF COGNITIVE IMPAIRMENT. *Psychiatria Danubina*, 34(suppl 1), 273-275.
- [6] Zhai, L. (2021). Ancient Literature and Art Based on Big Data. In 2020 International Conference on Data Processing Techniques and Applications for Cyber-Physical Systems: DPTA 2020 (pp. 357-365). Springer Singapore.
- [7] Laughlin, C. A. (2017). *The literature of leisure and Chinese modernity*. University of Hawaii press.
- [8] Cai, Z. Q., & Wu, S. (2019). Introduction: Emotion, patterning, and visuality in Chinese literary thought and beyond. *Journal of Chinese Literature and Culture*, 6(1), 1-14.
- [9] Mazanec, T. J. (2018). Networks of exchange poetry in late medieval china: Notes toward a dynamic history of tang literature. *Journal of Chinese Literature and Culture*, 5(2), 322-359.
- [10] Lin, S. (2020). The study of premodern Chinese literature in the digital era: new methods of quantitative statistics, databases, and visualization analyses. *library trends*, 69(1), 269-288.
- [11] Liang, H., Sun, X., Sun, Y., & Gao, Y. (2017). Text feature extraction based on deep learning: a review. *EURASIP journal on wireless communications and networking*, 2017, 1-12.
- [12] Bonch-Osmolovskaya, A., & Skorinkin, D. (2017). Text mining War and Peace: Automatic extraction of character traits from literary pieces. *Digital Scholarship in the Humanities*, 32(suppl_1), i17-i24.
- [13] Orekhov, B., & Fischer, F. (2020). Neural reading: Insights from the analysis of poetry generated by artificial neural networks. *Orbis Litterarum*, 75(5), 230-246.
- [14] Jindal Amar & Ghosh Rajib. (2022). Text line segmentation in indian ancient handwritten documents using faster R-CNN. *Multimedia Tools and Applications*, 82(7), 10703-10722.
- [15] Gabriel Antonesi, Tudor Cioara, Ionut Anghel, Ioannis Papias, Vasilis Michalakopoulos & Elissaios Sarmas. (2025). Hybrid transformer model with liquid neural networks and learnable encodings for buildings' energy forecasting. *Energy and AI*, 20, 100489-100489.