

Deep Learning-Based Analysis of Emotional Posture in Chinese Literary Works

Shengxiao Wang¹, 

¹ Zhengzhou College of Finance and Economics, Zhengzhou, Henan, 450000, China

ABSTRACT

In order to help readers analyze the emotions covered in Chinese literature and assist them in reading classic Chinese literature. In this paper, we propose a cross-granularity text sentiment analysis model based on interactive attention. Firstly, we carry out the preparation work for text processing of Chinese literary works, and for the shortcomings of the original word vector representation model that lacks the ability to distinguish text weights, we improve the TF-IDF algorithm in terms of weight imbalance of similar corpus, and combine the TF-IDF algorithm with the Word2vec word vector transformation model to enhance the weight of keyword vectors. A cross-granularity text sentiment analysis model based on Interactive Attention Network is proposed, in which the Interactive Attention Network is responsible for receiving the feature vectors encoded by RoBERTa and Word2vec respectively, capturing the bidirectional semantic information of long and difficult sentences. The sentiment analysis model in this paper has different degrees of accuracy and MF1 values on three different datasets. The method is able to provide a better sentiment visualization of some of the story lines of the classic Chinese literary work "Home", and it can accurately capture the emotional relationships in Chinese literature, which helps to assist readers in reading Chinese literary classics.

Keywords: Word2vec, TF-IDF, text sentiment analysis, interactive attention network, Chinese literature

1. Introduction

With the implementation of the strategy of "going out" of Chinese culture and the deepening of global cultural exchanges, the translation of Chinese literature, as an effective way of Chinese and foreign cultural exchanges and mutual understanding, especially in promoting traditional culture and enhancing the national cultural soft power, plays a crucial role [1-4]. However, there is a representative point of view that although a lot of efforts have been made to translate Chinese literature into foreign languages, the result is very little, the dissemination is poor, the acceptance is limited, and it is very difficult for Chinese literature to really go out or "go in" [5]. "Going in" is not to enter in the sense of

 Corresponding author.

E-mail address: wshxiao603@163.com (S. Wang).

Received 20 March 2024; Revised 25 May 2024; Accepted 10 December 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-243](https://doi.org/10.61091/jcmcc127b-243)

© 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/>).

physical space, but to understand, recognize and accept in the psychological sense, which is the most fundamental destination of literature going out, that is, “the purpose of going out is to go in” [6]. As an important part of the strategy of “culture going out”, translation of Chinese literature is a kind of cross-cultural communication activity, which not only aims at spreading Chinese culture and enhancing the image of the country in one way, but also aims at mutual understanding and mutual recognition in two ways [7-10].

In recent years, translation studies have successively noticed the absence of common readers' comments in the dissemination effect of translated texts, and have actively taken the path of empirical research to examine the dissemination and acceptance of translated texts among overseas common readers. Sentiment situation analysis refers to the process of analyzing, processing, inducting and reasoning about texts with subjective colors, which is of great significance in the fields of text mining, public opinion research and product word-of-mouth [11-14]. This technique can be effectively adapted to translation research, utilizing massive online review data to judge the overseas acceptance of translated texts [15-17]. By objectively reflecting the general overseas readers' emotional attitudes towards Chinese translated literature in a quantitative and qualitative way, it makes data-intensive literature overseas research possible to strengthen their identification with Chinese culture [18-20].

In the task of deep learning-based sentiment analysis of Chinese literature, this paper firstly carries out preprocessing work such as cleaning and word segmentation of text data, and then proposes the concept of word dispersion to improve the TF-IDF algorithm. Based on the sample characteristics of cross-granularity sentiment analysis, the pre-trained language model RoBERTa and word2vec word vector transformation model are used as the embedding layer of the cross-granularity text sentiment analysis model. The interactive attention mechanism is proposed to be embedded in the coding layer of the model, which improves the model's ability to learn the relevance of different text contents and its generalization ability. Finally, the performance evaluation of the cross-granularity sentiment analysis model and the evaluation of the effect of sentiment analysis are completed by the three public datasets and the dataset of Chinese literary classics, respectively.

2. Methods of Analyzing Emotional Situations in Chinese Literary Works

The basic process of text sentiment analysis using deep learning methods is shown in Figure 1, where the text data is preprocessed first, then the organized text is converted into word vectors, the constructed neural network model is used to extract the text features, and the final output of the text is the sentiment classification results. This chapter provides a systematic overview of the relevant theories and techniques for deep learning text sentiment analysis methods, including text preprocessing, text vectorization representation, deep learning techniques, pre-training language models and evaluation indexes.

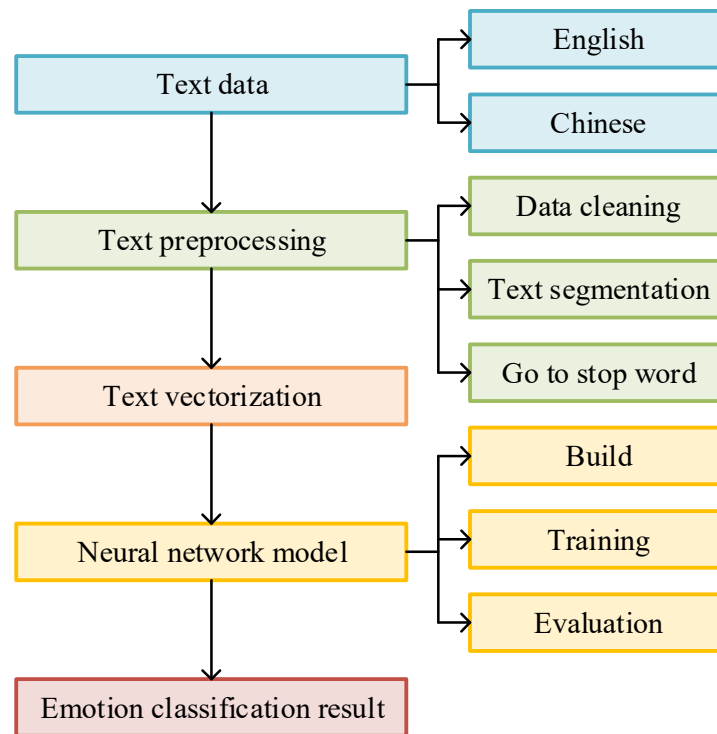


Fig. 1. Basic flow chart of sentiment analysis

2.1. Text pre-processing

When using deep learning methods for text categorization, large-scale data is required to support the training of neural networks. The current data obtained from different sources, especially social media, are usually unstructured, and the raw data may be interspersed with a large amount of noisy information. Therefore, it is necessary to preprocess the data before text analysis. Text preprocessing is to organize, filter and other necessary operations on the acquired raw data, so that the data can be more effectively utilized in the subsequent construction of the model. Generally speaking, text preprocessing includes data cleaning, text segmentation, deactivation, standardization and other operations.

1)Data Cleaning. Data cleaning is to organize the text data, clear the noise data, usually also use the regular matching of the data to match the rules. Data cleaning usually includes data format conversion, such as converting traditional Chinese characters to simplified Chinese characters; filtering of noise data, such as HTML tags, URL links, mailboxes, forwarding logos, and other data not related to sentiment analysis; and deleting or replacing emoticons in the text.

2)Text Lexicon. Text segmentation, as the name suggests, is to split text sentences into individual words. In this paper, Chinese and English datasets are used in the experiments, and the English and Chinese texts have different ways of word separation, English has a natural space as a separator, while Chinese word separation needs to be based on the semantic division of the word boundaries, and the two are different in terms of expression and syntactic structure. English text is relatively simple, can be separated by spaces, symbols, paragraphs as identifiers, through the regular expression to identify these punctuation marks to achieve the word division, to get a series of word groups. Different from English, Chinese does not have a natural delimiter such as space, Chinese punctuation is usually a pause, the tone of the expression of the role, and the Chinese expression of semantic basic unit is not the word, but the word composed of words, the length of the word is not uniform, so the need to use lexical methods to achieve the Chinese word segmentation. At present, Chinese word segmentation methods are mainly:

(1) Dictionary-based segmentation method. This method is in accordance with the dictionary of the text for rule matching, usually first build a Chinese thesaurus, the text and the content of the library to do a search to match the search to match the words in the thesaurus is recognized as a word, and its corresponding segmentation in the sentence.

(2) Segmentation method based on semantic understanding. This method utilizes semantic analysis and syntactic rules to segment the text and eliminate ambiguity.

(3) Statistical based segmentation method. This method does not rely on dictionaries and rules, but rather, the frequency of phrases is counted to split words, and the more times neighboring texts appear, the greater the chance of forming phrases.

(4) Deep learning based word separation method. Using deep learning technology, the neural network model is trained through a large number of datasets so that it has the ability to separate words and achieve the effect of word separation.

Based on Chinese participle technology, the participle method is integrated into the participle tool, and some universities and organizations have also launched participle tools, and Jieba participle tool is used for Chinese participle in this study.

3) Remove Deactivated Words. In the corpus data, there is a portion of words that do not convey specific semantic information, such as articles and prepositions like "a," "the," and "in" in English, and modal particles like "了," "啦," "的," and "吧" in Chinese. These types of words are referred to as stop words. In order to attenuate the noise effect of such words and improve the efficiency of model operation, deactivation word removal processing is required. Removal of deactivated words is the use of deactivated word lists to filter the text, i.e., after the text is segmented, the result is matched with the deactivated word list and deactivated words are eliminated.

4) Standardization. Standardization is mostly used to deal with English text data, including stemming and morphological reduction. Stem extraction is to remove the prefixes and suffixes of English words to get the root word, for example, "swimming", "swims" is reduced to "swim". Morphological reduction is to convert the complex form of an English word into its general form according to the dictionary, for example, to reduce "drove" to "drive". Standardization integrates words with different tense variations, which can reduce tense interference and the complexity of participles.

2.2. Text vectorization

2.2.1. Word2vec text representation. Word2vec [21] vector representation contains two important word vector composition methods, CBOW and skip-gram, both models use the input vocabulary to predict the surrounding text, but the overall idea and the scope of application differ greatly.

The CBOW word vector transformation method is mainly composed of input layer, projection layer and output layer. The method mainly sets the scale window and then utilizes the contextual vocabulary to train the current vocabulary according to the size of the window.

The objective function of CBOW word vector construction method is shown in Equation (1):

$$L = \sum_{w \in C} \log(p(w | \text{Context}(w))) \quad (1)$$

Where C denotes the total number of words in the dataset, w denotes the current word, and $\text{Context}(w)$ denotes the contextual word of the current word. The CBOW method word vector training process is as follows:

1) Input layer: input the contextual One-Hot encoded word vectors $v(\text{context}w1), v(\text{context}w2) \dots, v(\text{context}w2n)$, and set the dimension of the word vectors as K .

2) Projection layer: the input One-Hot encoding of the previous layer is summed and computed as shown in Equation (2):

$$x_w = \sum_{i=1}^{2n} v(contextw_i) \quad (2)$$

3) Output layer: each word constitutes a node, and the frequency of occurrence builds node weights, thus building a Huffman number structure used to reduce the amount of computation. The calculation formula is shown in (3):

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

where d_j^w represents the coding from the root node to the j th non-child node, d_j^w takes values ranging from 0 to 1, with 0 representing the left node and 1 representing the right node. x_w represents the vector of the projection layer, and θ_{j-1}^w represents the learning parameters of the j th non-child node. The formulae from the left and right sides through the Huffman tree are as follows:

$$\sigma(d_j^w = 1 | x_w \theta_{j-1}^w) = \frac{1}{1 + e^{-x_w^T \theta_{j-1}^w}} \quad (4)$$

$$\sigma(d_j^w = 0 | x_w \theta_{j-1}^w) = 1 - \sigma(d_j^w = 1 | x_w \theta_{j-1}^w) \quad (5)$$

The two are logarithmically summed and the objective function is obtained through the maximum likelihood formula as shown in equation (6):

$$L = \sum_{w \in c} \sum_{j=2}^r \left\{ (1 - d_j^w) \log \left(1 - \frac{1}{1 + e^{-x_w^T \theta_{j-1}^w}} \right) + d_j^w \log \left(\frac{1}{1 + e^{-x_w^T \theta_{j-1}^w}} \right) \right\} \quad (6)$$

Skip-gram, also known as the word skipping model, has a network model that is also divided into three layers, but the model is trained differently from the CBOW model. The model trains contextual words centered on the current word, just like the output of CBOW swaps places with the input.

The training process of Skip-gram is similar to that of CBOW, but the predicted probability $p(w | Context(w))$ is represented differently at the output layer, whose formula is shown in (7):

$$p(w | Context(w)) = \prod_{u \in Context(w)} p(u | w) \quad (7)$$

where the probability $p(u | w)$ that the context contains vocabulary u when the center word is w can be shown in equation (8):

$$p(u | w) = \prod_{j=2}^{l^u} p(d_j^u | x_w \theta_{j-1}^u) \quad (8)$$

Where d_j^u denotes the coding of the j th non-child node starting from the root node. d_j^u takes values in the range of 0 to 1, with 0 representing the left node and 1 representing the right node. x_w represents the vector of the projection layer, and θ_{j-1}^u represents the learning parameters of the j th non-child node.

After taking the logarithm, the objective function is obtained by the maximum likelihood formula as shown in equation (9):

$$L = \sum_{w \in c} \log \prod_{u \in Context(w)} \prod_{j=2}^{l^u} p(d_j^u | x_w \theta_{j-1}^u) \quad (9)$$

2.2.2. TF-IDF algorithm with improvements. In the field of sentiment analysis, the most frequent daily processing is text information, and the importance of different words in the text is expressed by the size of their weights, the larger the weight represents the stronger the influence of the word on the

text's sentiment tendency. TF-IDF [22] is a statistical method to assess the importance of different words for a given text data, commonly used in text mining and information retrieval.

TF denotes the word frequency, which when applied to sentiment analysis refers to the frequency of occurrence of a word in a given corpus, and its calculation formula is shown in Equation (10):

$$TF = \frac{n_{i,j}}{\sum_k n_{k,j}} \quad (10)$$

If a word occurs frequently, it is closely related to the theme of the text, but we also need to take into account the occurrence of some deactivated words that do not have practical significance, such as “the”, “had” and so on, the frequency is also very high, so the introduction of the inverse text frequency (IDF). IDF is a term that appears in many texts with a generally high frequency, indicating that the term is not representative of the current text and needs to be weakened. The formula for calculating IDF is shown in (11).

$$IDF = \log \frac{|D|}{|\{j : t_i \in d_j\}| + 1} \quad (11)$$

Where $|D|$ denotes the total number of documents in the dataset, and $\{j : t_i \in d_j\}$ denotes the number of documents in the dataset containing the current word t_i . The purpose of +1 is to avoid the denominator being 0. TF-IDF is obtained by multiplying TF and IDF with the following formula:

$$TF - IDF = TF * IDF \quad (12)$$

Through the operation of the above formulas, we finally get the TF-IDF value of the current word, which indicates the importance of the word to the text by the size of the value, which is helpful to improve the text classification effect to a certain extent.

In the traditional TF-IDF algorithm, TF simply considers that the higher the frequency of the feature word in the current article, the greater its importance. But this does not take into account the influence of its distribution on the weight. So in this paper, on the basis of TF, combined with the standard deviation formula, we propose a word dispersion D_i , which is calculated as follows:

$$D_i = \frac{\sqrt{\frac{1}{n-1} \sum_{i=1}^n ((tf_i(T_m) - \overline{tf(T_m)})^2)}}{\overline{tf(T_m)}} \quad (13)$$

where $tf_i(T_m)$ denotes the frequency of feature word T_m in the current class, $\overline{tf(T_m)}$ denotes the average frequency of T_m in all classes, and n denotes the total number of classes.

When T1 is evenly distributed in all classes, T2 is less distributed in a specific class, but more densely distributed in other classes, and there is a high variability of frequency counts among different classes. In the category analysis the feature term T2 is more representative than T1 for the recognition of specific categories, and it can also be calculated that the word dispersion of T2 is greater than that of T1, which means that the word dispersion has a proportional impact on the weights, and the new TF formula is obtained as follows:

$$TF = \frac{n_{i,j}}{\sum_k n_{k,j}} * D_i \quad (14)$$

In the traditional TF-IDF algorithm, the IDF is simply understood as the word text frequency is inversely proportional to the importance of the word, but for most of the corpus information, this is not quite correct. The IDF structure is relatively simple, and it can't extract the keyword information completely and effectively. In order to solve the above problems, this paper replaces the ratio of the total number of documents to the number of documents containing the current word with the ratio of

the total number of words in the corpus to the number of times the current word to be checked appears in the corpus. This treatment reduces the influence of similar texts in the corpus on the word weight, and expresses the importance of the current word in the text to be checked more effectively. The formula is as follows:

$$IDF = \log \frac{\sum_{i=1}^m nt_i}{nt_i} \quad (15)$$

where nt_i denotes the total number of occurrences of feature word t_i in the corpus. The final improved TF-IDF formula is obtained as follows:

$$TF-IDF = \frac{n_{i,j}}{\sum_k n_{k,j}} * D_i \times \log \frac{\sum_{i=1}^m nt_i}{nt_i} \quad (16)$$

2.2.3. Improved word2vec text vector representation. The collected text dataset is transformed into text vectors in order to effectively extract its internal features. Word2vec model is a very important tool for vector transformation, which mainly contains two word vector composition methods, skip-Gram and CBOW, both of which use the input vocabulary to predict the surrounding text.

In this paper, the Word2vec model and the TF-IDF algorithm both complementary methods are combined to obtain weighted word vectors. In this paper, the word vectors trained by Word2vec model are combined with the improved TF-IDF for weight calculation. The improved TF-IDF formula (17) is utilized to compute the weight value of the word in the sentence $K(t, D_i)$, where $D_i (i=1, 2, \dots, M)$ denotes the set D containing M words.

$$K(t, D_i) = \frac{TF(t, D_i) \times IDF(t)}{\sqrt{\sum_{t \in D_i} [TF(t, D_i) \times IDF(t)]^2}} \quad (17)$$

where the denominator is the normalization factor.

For each sentence $D_i (i=1, 2, \dots, M)$, where the sentence vector can be expressed in the following form:

$$d_i = \sum_{t \in D_i} w_t K(t, D_i) \quad (18)$$

In this way, we get the sentence vector d contains both the interrelated features of the text, and also the improved TF-IDF weight matching of the key information of the sentence, so that the sentence vector is more comprehensive in feature extraction of the text, which is more conducive to the next work of model classification.

2.3. Cross-granularity text sentiment analysis model based on interactive attention

The cross-granularity text sentiment analysis model proposed in this paper consists of: an embedding layer that independently models the two texts, an encoding layer, two structurally consistent attention networks, and a sentiment polarity classifier.

The cross-granularity text sentiment analysis method proposed in this chapter applies the Transformer [23] encoder structure and has the robust property of pre-training model RoBERTa to extract the semantic features and sentence structure information of the document-level review text. At this stage, the model encodes the words in the review text sequence by BPE character encoding, and fuses the features with paragraph encoding and relative position encoding.

$$E_{word} = E_{token} \oplus E_{pos} \oplus E_{seg} \quad (19)$$

$$E_{word} = (O_{token} \quad O_{pos} \quad O_{seg}) \begin{pmatrix} W_{token}^{V \times H} \\ W_{pos}^{P \times H} \\ W_{seg}^{S \times H} \end{pmatrix} \quad (20)$$

Eqs. O_{token} and O_{pos} are one-hot encoding vectors of subwords encoded in BPE characters with relative position information, respectively, in this paper, the method is to input the complete review text into the RoBERTa embedding layer, and O_{seg} all-zero vectors denote the real text that is not filled in the input sequence. The embedding matrices E_{oken} , E_{pos} and E_{scg} are obtained by multiplying the three one-hot vectors and the parameter matrices $W_{token}^{V \times h}$, $W_{pos}^{P \times h}$ and $W_{seg}^{S \times h}$, and V, P, and S are the dictionary sequence dimension, the maximum value of the coding position, and the input sentence attribute value, respectively, and h is the hidden layer dimension, which is set to 768 dimensions in the RoBERTa model. E_{wond} is the model input matrix containing contextual correlation information, and symbol $\oplus$ denotes the overlay operation against the embedding matrix. The RoBERTa model is computed with 12 layers of multi-head self-attention to capture long-distance word dependencies in long text cross-sentence comments, and the high-quality feature engineering is the basic guarantee for the performance of the whole sentiment analysis model.

The summary part of the cross-granularity samples mainly consists of exclamatory sentences, imperative sentences or several phrases that express the customer's feelings, with a clear and concise sentence structure and without complex contextual relationships. In order to control the complexity of the model, we choose to add a static word vector word2vec tool to the embedding layer, which is responsible for portraying wordiness and inter-word relationships in the summary sequence. The embedding layer of the cross-granularity text sentiment analysis model consists of two parallel-structured language models, RoBERTa and Word2Vec, which ultimately output two feature matrices with different dimensions.

The cross-granularity sentiment analysis model proposed in this paper uses an LSTM network to build a coding layer that passes the feature vectors output from the RoBERTa model and word embedding matrix to the interactive attention network. LSTM is a special kind of recurrent neural network, characterized by getting rid of the simple RNN memory superposition of the way of working, with more powerful memory capacity can be applied in the intention recognition, machine translation, stock prediction and other tasks to learn the sequence relationship.

Inside the LSTM [24] cell, there are cell states C_t that preserve historical information and hidden states h_t that are passed over time steps, and the interaction of information inside the cell is controlled by forgetting gates f_t , input gates i_t , and output gates o_t . The LSTM network takes the coded embedding layer output eigenvectors x_t as the input signals, and realizes the computational process of storing and transferring the information as follows, with W and b denoting the weight matrix and bias parameter vectors of the stages.

$$f_t = \sigma(W_f \cdot [C_{t-1}, h_{t-1}, x_t] + b_f) \quad (21)$$

$$i_t = \sigma(W_i \cdot [C_{t-1}, h_{t-1}, x_t] + b_i) \quad (22)$$

$$o_t = \sigma(W_o \cdot [C_{t-1}, h_{t-1}, x_t] + b_o) \quad (23)$$

$$\tilde{C}_t = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c) \quad (24)$$

$$C_t = f_t * C_{t-1} + i_t * \tilde{C}_t \quad (25)$$

$$h_t = o_t * \tanh(C_t) \quad (26)$$

The RoBERTa submodel in the embedding layer outputs document-level review text word vectors, and the LSTM encoding network obtains the long-sentence review text hidden features $[h_1^r, h_2^r, \dots, h_n^r]$. The submodel applying the Skip-grams algorithm within the embedding layer does not have a bidirectional noise-reducing encoder structure, and it cannot adequately incorporate the contextual knowledge in the word vectors. In order to allow the model to encode the word representations formed by summary conversion from the perspective of the order relationship, the design embeds this part of the word into the bidirectional structure of the LSTM network within the input coding layer, with the aim of giving the model stronger feature fitting ability by deepening the depth of the network. BiLSTM accumulates both forward and reverse serialized dependency information, and splices them together to obtain a combination of positive and negative order, which comprehensively summarizes the Sentence information ground hidden state $[h_1^s, h_2^s, \dots, h_m^s]$, the i rd word of the sequence is denoted as $h_i^s \in R^{2d}$. Here d represents the number of neurons of the LSTM network.

The key feature of the cross-document sentiment analysis model is the introduction of an attention network that facilitates its mutual analysis by exploiting the semantic relationships between REVIEW TEXT and SUMMARY within the cross-granularity text samples. The core layer is cross-computed in a way that allows the model to learn the background knowledge contained in another part of the review. The output states of the two coding layers of the model are pooled to obtain the average vectors r_{avg} and s_{avg} in the time-step dimension, which are regarded as the initial representations of the REVIEW TEXT and SUMMARY input to the interactive attention layer, and the process of calculating the probability of the attention distributions a_i, β_i is described in Eqs. (27) and (28).

$$a_i = \frac{\exp(\gamma(h_i^r, s_{avg}))}{\sum_{j=1}^n \exp(\gamma(h_j^r, s_{avg}))} \quad (27)$$

$$\beta_i = \frac{\exp(\gamma(h_i^s, r_{avg}))}{\sum_{j=1}^m \exp(\gamma(h_j^s, r_{avg}))} \quad (28)$$

The γ in the expression is the importance scoring function computed to reflect the degree of correlation between the review text via the coding layer output h_i^r and the mean vector s_{avg} that characterizes the summary information. The correlation between h_i^s and r_{avg} is calculated in the same way. Mean pooling aggregates all the time-step features in the sequence into vectors and induces the interactive attention mechanism to generate reasonable weight coefficients, where W_r, b_r and W_s, b_s represent the weight parameters and biases to be learned within functions $\gamma(h_i^r, s_{avg})$ and $\gamma(h_i^s, r_{avg})$, respectively. The algorithm sets up this step to allow the model to learn the clues provided by another part of the sample when analyzing the text in a certain place.

$$\gamma(h_i^r, s_{avg}) = \tanh(h_i^r \cdot W_r \cdot s_{avg}^T + b_r) \quad (29)$$

$$\gamma(h_i^s, r_{avg}) = \tanh(h_i^s \cdot W_s \cdot r_{avg}^T + b_s) \quad (30)$$

The attention weight a_i, β_i is computed and weighted and summed with the output states of the encoding layer, thus obtaining the feature vector r, s that implies the semantic information of review text and summary, which are used as the two structurally coherent and informationally interactive attention outputs, respectively.

$$r = \sum_{i=1}^n a_i h_i^r \quad (31)$$

$$s = \sum_{i=1}^m \beta_i h_i^s \quad (32)$$

Splice r and s by the dimension of the hidden layer to get the vector d that integrates the deep features of review text and summary. The fully connected layer at the top of the model constructs a nonlinear mapping by parameter matrix W_g and bias b_g , whose function is to integrate the sentiment knowledge and abstract semantic information of the sample sequences, which is finally converted into the input vector k for the classifiers.

$$k = \tanh(W_g \cdot d + b_g) \quad (33)$$

The feature vector k obtained from the conversion of the two text hidden features by hyperbolic tangent function is input into the top layer classifier of the cross-granularity text sentiment analysis model, which is used as the basis for the discrimination of the final category labels. Softmax classifier sets the label with the highest probability in the total number of sentiment categories C to be the result of the sentiment analysis. Here W and b are the weight parameters and bias of the classification layer function, and the model completes the sequence to category conversion in this step, and the output value $y_i (i \in [1, C])$ represents the predicted label of the cross-granularity text sentiment analysis model.

$$y_i = \text{softmax}(W \cdot k + b) \quad (34)$$

3. Results of model performance analysis

3.1. Data sets

This chapter continues to use the two public datasets released by the International Assessment Competition in 2014; Restaurant, Laptop, and the Twitter social media dataset for experimental validation. The dataset details are shown in Table 1.

Table 1. Three public data collection label statistics

Data set	Positive emotion	Neutral emotion	Negative emotion
Restaurant-Train	2169	636	808
Restaurant-Test	737	197	194
Laptop-Train	991	466	865
Laptop-Test	338	176	125
Twitter-Train	1573	3126	1553
Twitter-Test	178	346	170

3.2. Comparative Experimental Results

In order to test the validity of the models in this paper, this paper will be compared with LSTM model, ATAE-LSTM model, IAN model, RAM model, AEM model, ASGCN model, IGATs model, GL-GCN model, DBG-GBGCN model, KDGC model, KDGC model, HGCN model, SSEGCN model, SK-GCN1+ BERT model, SK-GCN2+BERT model, R-GAT+BERT model, and CGAT+BERT model for comparative analysis.

This section compares the results of this paper's model and the comparison model on the aspect-level sentiment analysis task on the three datasets. The experimental results are shown in Table 2, where the bolded font indicates the optimal results in the current metrics. Compared with the other proposed comparison models, the model proposed in this paper improves Acc by 0.23%, 0.36%, and 0.81% on the Restaurant, Laptop, and Twitter datasets, respectively, and MF1 improves MF1 by 1.53%, 0.74%, and 1.52%, respectively. It shows that the design method of the model in this paper can well help the model to better capture the sentiment dependency between aspect and context, which makes the model performance improved.

Table 2. Results in three public data sets

Model	Restaurant		Laptop		Twitter	
	Acc	MF1	Acc	MF1	Acc	MF1
LSTM	76.67	64.73	67.91	61.58	67.32	65.02
ATAE-LSTM	77.12	67.27	68.42	63.7	69.7	67.62
IAN	78.88	70.64	71.81	67.77	72.01	70.08
RAM	80.22	70.53	73.98	70.72	69.89	67.88
AEN	81.04	72.16	73.52	69.05	72.87	69.8
ASGCN	80.81	72	75.59	71.03	72.12	70.39
IGATs	82.33	74.07	76.03	72.15	75.32	73.39
GL-GCN	82.06	73.48	76.91	72.75	73.25	71.24
DBG-GBGCN	82.3	74.21	76.78	72.68	73.97	72.62
KDGN	87.01	82.01	81.28	77.57	77.64	75.55
HGCN	87.46	82.18	81.49	77.36	77.5	76.45
SSEGCN	87.34	81.07	80.99	77.97	77.43	76.06
SK-GCN1+BERT	81.87	73.43	79.28	75.08	74.72	73.38
SK-GCN2+BERT	83.41	75.22	78.99	75.58	75.02	72.94
R-GAT+BERT	85.06	79.26	79.04	75.74	75.14	74.01
CGAT+BERT	86.23	80.37	80.41	76.49	77.44	76.34
DC-GCN	88.13	83.25	81.77	77.1	77.49	76.79
This model	88.33	83.55	82.43	78.28	78.06	77.96

3.3. Attention visualization results

To more intuitively show how the proposed model in this paper improves the performance of predicting aspect-based sentiment polarity, this subsection visualizes the attention weights through a one-sided example and a multifaceted example. Its exemplary one-sided example and multifaceted example attention distribution graphs are shown in Figures 2 and 3, respectively. The proposed model in this paper pays more attention to key emotion words in extracting emotion features corresponding to specific aspects. In particular, the multifaceted examples show that the proposed model can distinguish different aspects and focus on different contextual sentiment words according to different aspects. This suggests that the improved word2vec text extraction method utilized to extract aspect words in the opinion span range and the introduction of a syntactic adjacency matrix enhanced with external sentiment knowledge can accurately capture aspect-specific sentiment features in the opinion span range and improve the performance of aspect-level sentiment analysis.

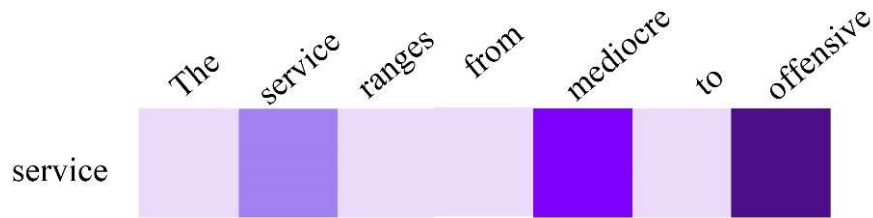


Fig. 2. Sample diagram of a single aspect sentence

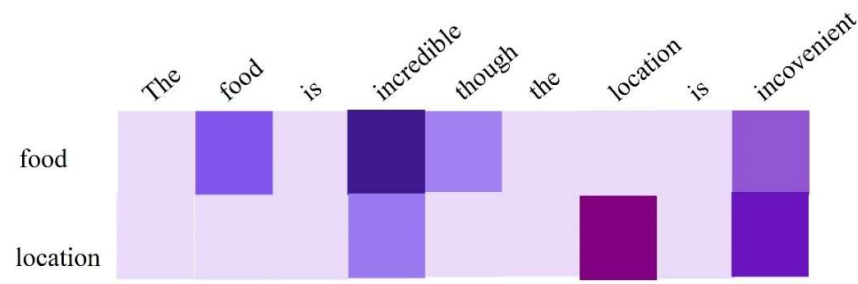


Fig. 3. A sample diagram of a sentence in multiple aspects

3.4. Optimal parameter selection for the model

Fig. 4 and Fig. 5 show the test accuracy under different embedding dimensions and different sliding window sizes, respectively. When the embedding dimension is 30 and the sliding window size is 10, the test accuracy is the highest, which is 0.839 and 0.837, respectively. Therefore, the embedding dimension and sliding window size are set to 30 and 10, respectively.

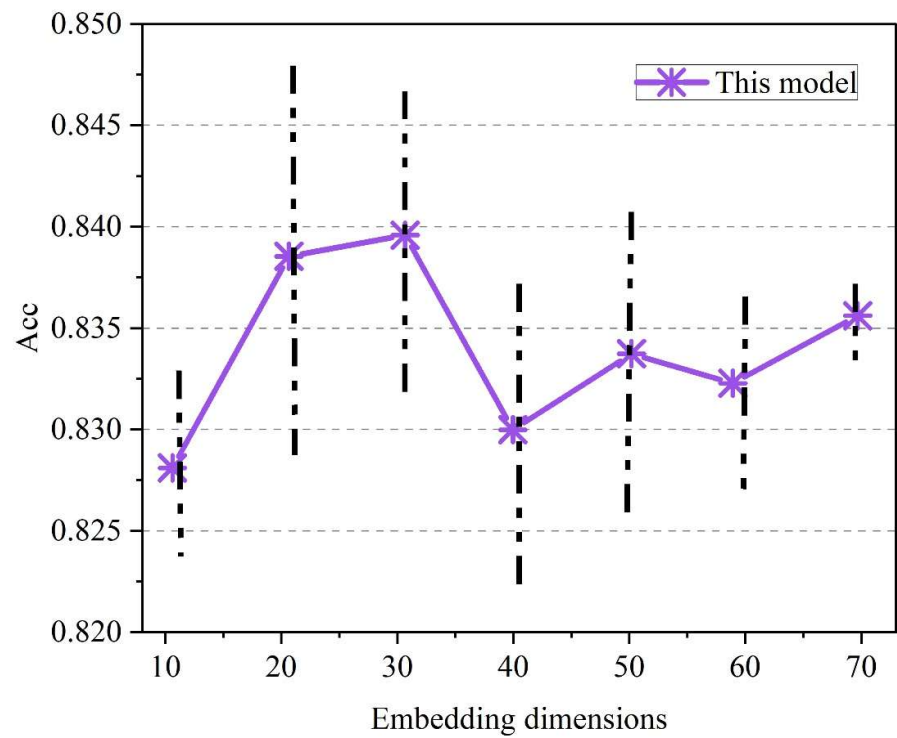


Fig. 4. Test accuracy of different embedding dimensions

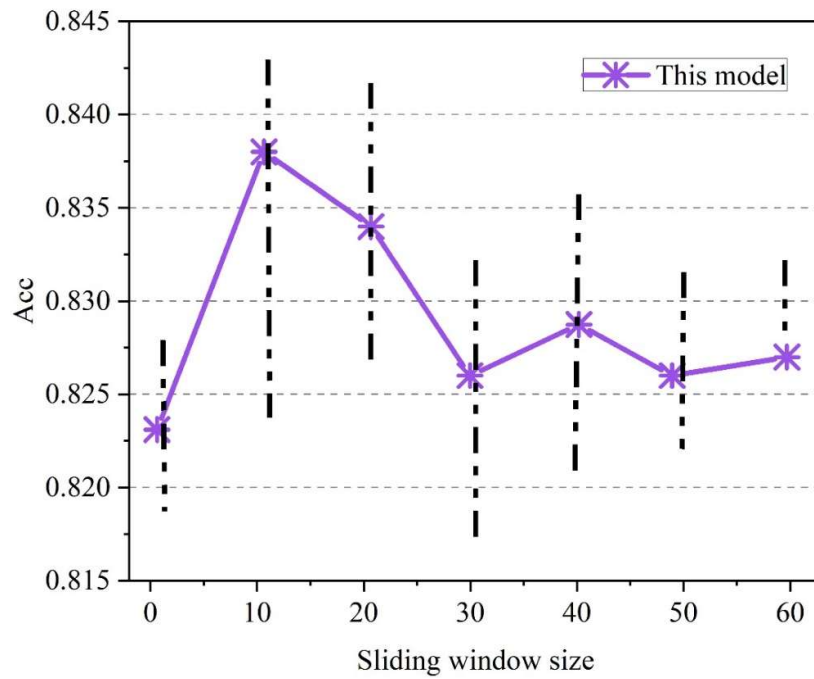


Fig. 5. Test accuracy of different sliding window sizes

4. Modeling Sentimental Posture Analysis

4.1. Analysis of visualization results

Visualizing the emotional routes of a large number of stories summarizes six core types of emotional arcs, and Figure 6 shows one of these emotional routes, and the visualization of the emotional routes of the 20 books closest to this arc. The visualization effect of this method is more cluttered, and the reader cannot understand the general plot of the story based on the chart at all.

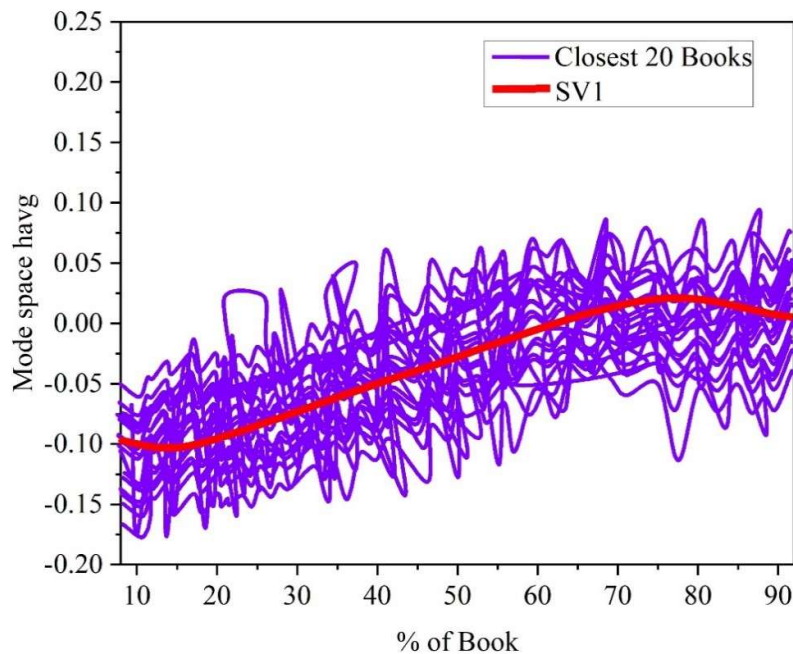


Fig. 6. Emotional arc (one of six)

In this paper, we add a rich visual information coding design based on the results of character emotion categorization to the narrative story line generated by the iStoryline tool. The aggregation of curves represents that in the current time period, the character represented by the line is in the same scene at that time period. Figure 7 shows the visualization of the story line in chapters 19-25 of *Home*. The figure clearly shows the direction of the story: the younger members of the Gao family go to the water pavilion to play, Shuzhen is afraid of ghosts and haunts Mingfeng, and the brothers Juehui and Jiemin play with them. The brothers Jiehui and Jiemin play with them. They sit in a boat and sing and play flutes in the lake, and then there is a heated discussion on the issue of foot-binding, marriage, and schooling. At this time, news comes that Zhang's army is coming in, and everyone is so frightened that Cousin Mei, Qin and Mrs. Zhang all come to take refuge. As the war became more and more tense and the Gao family was the richest in the area, people from the third, fourth and fifth houses went to their relatives' houses for refuge. After the war, all is back to peace. Ruijue and Mei are instantly attracted to each other, and the two of them talk openly, while the relationship between Jiemin and Qin also heats up quickly. The plot of the whole story is very clear, and the story can be clearly understood according to the chart. It shows that the accuracy of this paper's sentiment analysis model in categorizing sentiment is more excellent on the dataset of Chinese literary works. The emotion of each character in the novel is categorized. Adding emotional elements makes the automatically generated story line more vivid, full and beautiful.

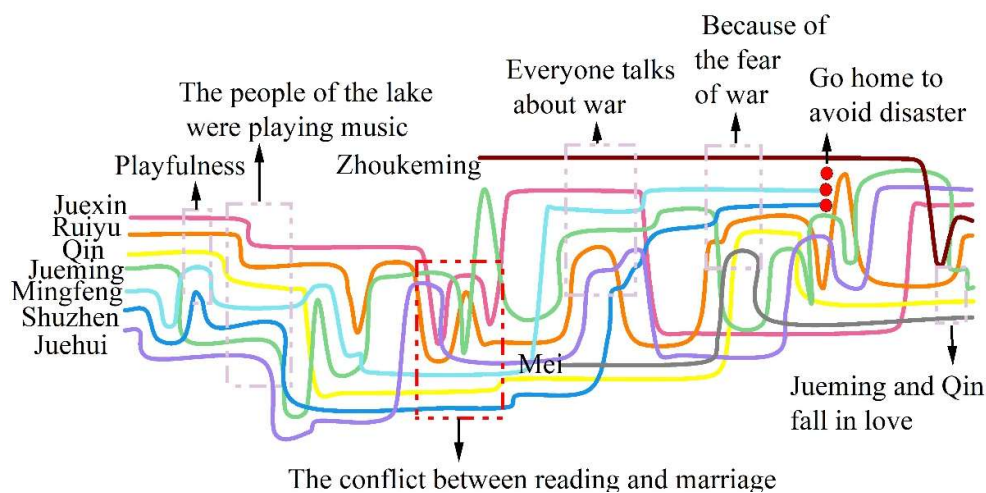


Fig. 7. Home story line visualization

4.2. Emotional Vocabulary Analysis

The study selects 600 poems describing the landscape of Gyeonggi, screened from Gu Taiqing Yiyi Poetry Collection, for text structuring. The nouns of different types of landscapes were sorted by word frequency according to five types: architecture, plants, mountains, water systems, and others, and finally the top thirteen words of each type of landscape were screened and summarized, and the summary results are shown in Table 3. From the table, it can be seen that the most frequent words in the high-frequency words are the meteorological word "east wind", and even many wind-related words such as "wind blowing" and "cool breeze" occupy a high position in the word frequency order. It can be seen that the climate of Gyeonggi is of great significance to the shaping of the literary landscape.

Table 3. Summary of high frequency landscape words in Jingji Ji Tourism poetry

Architecture		Plants		Mountain		Drainage pattern		Other	
Biwa	10	Haitang	29	Xishan	26	Liushui	23	Dongfeng	38
Changfang	10	Taohua	28	Nanggu	23	Qiushui	10	Fengchui	23
Shuangqiao	10	Qinqin	22	Shangzhong	15	Cixi	9	Liangfeng	19
Xiaolou	9	Yangliu	15	Yilu	14	Chunshui	9	Jinnian	19
Renjia	9	Huanghua	15	Yuanshan	14	Qushui	9	Xieyan	18
Tianyou	9	Huakai	13	Gaofeng	15	Jingshui	9	Shili	16
Simian	8	Caomu	12	Nangshan	10	Taipinghu	8	Xifeng	15
Qinfeng	8	Laoshu	11	Shanshui	10	Hanquan	8	Bainian	15
Xuting	8	Lihua	10	Shantou	9	Kongshui	7	Linlong	14
Shanglou	8	Cangtai	10	Gudong	9	Kunminghu	7	Fengyu	13
Tingyuan	8	Xiuzhu	10	Shengshan	9	Yeshui	6	Qiufeng	12
Langan	7	Hehua	10	Qinshan	8	Hushang	6	Chunfeng	11
Nanguo	7	Luohua	9	Denggao	8	Xishui	5	Chengnan	11

Ancient Chinese literati often give subjective emotions to objective scenery to establish imagery, and in literary creation to show the subjectivity of the landscape. In China's traditional landscape aesthetics, love and scenery mingling and interpenetration, the scene is all love, love figurative become the scene, the landscape landscape emotional analysis is also a special interpretation of the landscape ontology, JiYou poets emotional expression and poetic transmission of the base are mainly natural landscape. Secondly, the landscape elements are combined according to different sequences, and the different atmospheres created are the paths for the formation of emotional colors in the landscape space. Therefore, the emotional change is one of the important aspects of the landscape spatial sequence to be discussed, from which the emotional concentration of the poems can be judged as a reference for the recovery of the spatial sequence of the poetic description. The emotions expressed in Jiyou's poems and literary works reflect the evaluation and feelings of natural geography and landscape environment and cultural space. Thus, although literary emotions do not have a tangible form, they are often closely related to the spatial elements of the landscape. The emotional factors of different individuals in the poems of Gu Taiqing and Yi Ei give the Gyeonggi landscapes more vivid imaginative space, so it is necessary to explore the emotional tendency of the Gyeonggi landscapes from the emotional point of view, and to get the evaluation and feeling of the landscape space of Gyeonggi.

Through the cross-granularity text sentiment analysis model based on interactive attention designed in this paper for sentiment analysis, the emotional tendencies contained in the textual content of Gyeonggi landscape poems were quantitatively expressed to obtain the emotional evaluation of Gu Taiqing and Yi Yi's natural and humanistic landscapes in the Gyeonggi region. The emotions are categorized into positive, neutral, and negative, and the specific emotion percentages are shown in Table 4, in which positive emotions account for the highest percentage of 48.25% in the emotion categories, followed by negative emotions accounting for 31.39%, and finally neutral emotions accounting for 20.36%. First of all, jiyou poems are mainly the products of the ancient literati's travel and appreciation activities, and most of the time, the traveling behavior is based on the relaxed and pleasant emotions, coupled with the human nature's closeness to nature, the jiyou poems tend to have positive emotional preferences in their depiction of the landscape, so it is understandable that the positive emotions account for the highest percentage in the overall emotional analysis of the jiyou poems. Secondly, the negative emotions are mainly extracted from the specific descriptions of bad weather, dangerous terrain, and desolate environment in the poems, as well as from the general negative imagery expression generated by special time and special landscape, such as the feeling of grief and sunset conveyed by the clearing of the sky and dusk, and the solemn and mournful

atmosphere brought by the temples and tombs. In the end, the proportion of negative emotions in the overall emotional analysis is much lower than that of positive emotions, which leads to the conclusion that the overall emotional evaluation of the Gyeonggi landscape space by Gu Taiqing and Yi Ei is positive, and the feelings obtained based on the Gyeonggi landscape space are pleasant.

Table 4. Analysis of landscape emotion in Jingji Province

Emotions	Proportion/%	Strength	Proportion/%
Positive	48.25	General	38.33
		Medium	7.24
		Altitude	2.68
Neutrality	20.36	/	/
Negativity	31.39	General	25.89
		Medium	4.52
		Altitude	0.98

5. Conclusion

In this paper, we improved the existing models for sentiment analysis and designed a cross-granularity text sentiment analysis model based on interactive attention, and the summary is summarized as follows:

The sentiment analysis model in this paper obtains optimal results on three different datasets. The cross-granularity text sentiment analysis model based on interactive attention improves the accuracy on Restaurant, Laptop and Twitter datasets by 0.23%, 0.36%, 0.81%, and the MF1 by 1.53%, 0.74%, 1.52%, respectively. It indicates that the attention mechanism added to the model in this paper can better help the model capture sentiment dependencies and improve the model's sentiment classification effect.

In terms of the model's visualization effect on sentiment analysis, this paper's model visualizes the story line of chapters 19-25 of the classic Chinese literary work "Home". The visualization results clearly show the direction of the whole story, and the behavior and emotional posture of each character are presented completely within a certain timeline. It shows that the cross-granularity text sentiment analysis model based on interactive attention is able to accurately classify the characters' emotions in Chinese literature. Adding the obtained emotion classification results to the story line presents a more vivid, full and beautiful story, which positively promotes the assisted reading of Chinese literary works.

References

- [1] Xu, J. (2024). On the Outbound Translation Behavior and Criticism in the Dissemination of Chinese Classics: Review of A Study on Yang Xianyi's Translation. *New Thoughts on Translation*, 195-207.
- [2] McDougall, B. S. (2014). World literature, global culture and contemporary Chinese literature in translation. *International Communication of Chinese Culture*, 1, 47-64.
- [3] Jiang, M. (2019). Theoretical Proposal for Studies on the Outward Translation of Contemporary Chinese Literature: An Interdisciplinary and Eclectic Approach. *Translation Quarterly*, (93).
- [4] Geng, Q. (2022). Discourse Models on Practice and Studies of Outbound Translation of Chinese Literature. In *Chinese Literature in the World: Dissemination and Translation Practices* (pp. 39-56). Singapore: Springer Nature Singapore.
- [5] Xu, J. (2024). Theoretical Reflection and Research Approach of Studying the Translation of Chinese Literature. In *New Thoughts on Translation* (pp. 183-193). Singapore: Springer Nature Singapore.

- [6] Cao, S., GAO, J., & JIANG, Y. (2018). On Wang Rongpei's Drama Translation Strategy: A Case Study of The Peony Pavilion. *Studies in Literature and Language*, 17(3), 28-34.
- [7] Guo, J., & Zou, D. (2023). Reception study: The omission of narrative text in the English translation of Mo Yan's *Life and Death Are Wearing Me Out*. *Frontiers in Communication*, 8, 1063490.
- [8] Wang, B. (2022). Translating Chinese Literature to Reach an Audience in China: Lin Yijin and Lu Xun's Stories in English. In *Chinese Literature in the World: Dissemination and Translation Practices* (pp. 139-158). Singapore: Springer Nature Singapore.
- [9] Tor-Carroggio, I., & Rovira-Esteva, S. (2021). Chinese literary translation in Spain up until 2020: A quantitative approach of the who, what, when and how. *Skase: Journal of Translation and Interpretation*, 14(1), 67-94.
- [10] Wu, Y., & Jiang, M. (2024). Revisiting directionality in China's "outward translation". *International Journal of Applied Linguistics*.
- [11] Zhao, L., & Luo, B. (2024). The Reception of Giles' and Minford's English Translation of *Liaozhai Zhiyi*: A Sentiment Analysis Perspective. *International Journal of Chinese and English Translation & Interpreting*.
- [12] Jidong, Z. H. A. N. G., & Ruofan, D. U. (2024). A Comparative Study of Sentiment Words in Literary Translation A Case Study of *Who Do You Think You Are* and Its Chinese Translations. *Journal of Foreign Languages*, 47(2), 117-128.
- [13] Wu, K. (2021). Reception of Jin Yong's Wuxia Novels in English and French: A Sentiment Analysis with Reflection. *International Journal of Linguistics, Literature and Translation*, 4(3), 117-123.
- [14] Hou, Y., & Ma, Y. (2023, October). Cross-Lingual Sentiment Analysis in Literary Translation: A Case Study of the Novel *Crystal Boys*. In *2023 10th International Conference on Behavioural and Social Computing (BESC)* (pp. 1-6). IEEE.
- [15] LIU, W., & ZOU, J. (2023). A Study on the Acceptance of the English Translation of *Romance of the Three Kingdoms* by Overseas Readers Based on Python Data Analysis Technology. *Sino-US English Teaching*, 20(5), 187-194.
- [16] Wang, F., Humblé, P., & Chen, W. (2019). A bibliometric analysis of translation studies: A case study on Chinese canon the journey to the West. *Sage Open*, 9(4), 2158244019894268.
- [17] Yang, L., & Zhou, G. (2022). A semantic similarity analysis of multiple English translations of *The Analects*: Based on a natural language processing algorithm. *Frontiers in Psychology*, 13, 992890.
- [18] Abdullah, N. A. S., & Rusli, N. I. A. (2021). Multilingual Sentiment Analysis: A Systematic Literature Review. *Pertanika Journal of Science & Technology*, 29(1).
- [19] Menghan, Q. (2017). Penguin classics and the canonization of Chinese literature in English translation. *Translation and Literature*, 26(3), 295-316.
- [20] Wu, K., & Li, D. (2022). Are translated Chinese Wuxia fiction and western heroic literature similar? A stylometric analysis based on stylistic panoramas. *Digital Scholarship in the Humanities*, 37(4), 1376-1393.
- [21] Arar Al Tawil, Laiali Almazaydeh, Doaa Qawasmeh, Baraah Qawasmeh, Mohammad Alshinwan & Khaled Elleithy. (2024). Comparative Analysis of Machine Learning Algorithms for Email Phishing Detection Using TF-IDF, Word2Vec, and BERT. *Computers, Materials & Continua*(2), 3395-3412.
- [22] Yonglian Luo & Cailin Lu. (2024). TF-IDF combined rank factor Naive Bayesian algorithm for intelligent language classification recommendation systems. *Systems and Soft Computing* 200136-200136.
- [23] Qin Jiang, Qinglin Wang, Lihua Chi & Jie Liu. (2025). Cross-scale hierarchical spatio-temporal transformer

for video enhancement. Knowledge-Based Systems112773-112773.

- [24] Chunzi Wang,Fusheng Xie,Junpeng Yan & Yiqing Xia. (2025). A U-MIDAS modeling framework for forecasting carbon dioxide emissions based on LSTM network and LASSO regression. Energy Reports16-26.