

Research on Innovative Application of English Business Writing Intelligent Assisting System Based on Image Recognition Algorithm and Language Modeling

Zhenying Zhang¹, 

¹ Department of Basic Courses, Shangqiu Institute of Technology, Shangqiu, Henan, 476000, China

ABSTRACT

This paper focuses on the demand for intelligent assistance in English business writing scenarios and proposes an intelligent assistance system for English business writing based on image recognition algorithm and language model. The system is able to quickly extract image information related to the writing topic through the similarity vocabulary matching technology combined with the image retrieval recognition function based on CBLSTM-Attention model. The language model is utilized to make accurate vocabulary recommendation and expression for the writing scene and user input content, and finally construct the overall framework of the intelligent assistive system based on English business writing. The system performs well in terms of vocabulary matching accuracy and writing efficiency improvement, with an average matching accuracy of over 90%. Students' quality of writing is essentially improved with the help of the system in this paper. The actual case study shows that studying under the intelligent assistance system, the post-test scores of English business composition of the students in the experimental class increased significantly by 9.9393 points ($P < 0.05$) compared with the average scores of the control class, and it is obvious that applying the model of this paper to the classroom teaching can lead to a significant improvement in the performance of the students, which demonstrates the good prospect of its application.

Keywords: Vocabulary Matching Technique, Image Retrieval, CBLSTM-Attention Model, Language Modeling, English Writing

1. Introduction

Business English writing is a skill of using English to communicate in writing in business scenarios, and it plays an important role in international business communication [1-2]. The basic principles of business English lie in the aspects of clarity, appropriateness, simplicity, accuracy, professionalism,

 Corresponding author.

E-mail address: 1350013064@qqxy.edu.cn (Z. Zhang).

Received 08 February 2024; Revised 18 April 2024; Accepted 30 November 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-303](https://doi.org/10.61091/jcmcc127b-303)

© 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/>).

reliability, time management and attention to language style, etc. Only by paying attention to these principles, can we conduct business English writing better and provide high-quality business documents and communication [3-6]. While traditional business English writing is mainly written manually, there are deficiencies in cognitive deficiencies, subjective emotions, etc., with the development of artificial intelligence, the birth of intelligent writing assistance system avoids the deficiencies of manual writing and makes business English more perfect [7-10].

Intelligent writing assistance system is a kind of application based on artificial intelligence technology, which uses image recognition algorithm and language modeling aimed at helping users to improve writing efficiency and writing quality, which has a wide range of application value in the fields of academics, business and creative writing, and has a positive effect on the improvement of users' writing ability [11-14]. And in business English writing, the intelligent writing assistance system, after understanding the user input, is able to generate corresponding writing suggestions according to the user's writing purpose and content, which includes suggestions for grammatical correction, vocabulary substitution, sentence restructuring, etc., which improves the efficiency of business English writing and the quality of writing, and provides the user with personalized and precise writing assistance [15-18]. With the continuous development and improvement of artificial intelligence technology, the intelligent writing assistance system is expected to develop into an intelligent and efficient writing tool in the future, further promoting people's writing ability and creativity [19-21].

Literature [22] proposed an intelligent support system for English writing based on the B/S model, which is based on a corpus of vocabulary, grammar rules, etc., and utilizes the natural language processing technology, which combines rule matching and probability statistics to evaluate and optimize the efficiency of writing, revealing that the system can effectively improve the teaching direction. Literature [23] describes the application of ChatGPT in writing English writing, ChatGPT is able to answer all the questions on various topics, such as writing in English, telling personal experiences, etc., and takes into account the order of writing, morphology, use of tenses, etc., but the grammatical accuracy it exhibits needs to be studied in depth. Literature [24] built an AI-based English writing system based on machine learning and flock effect algorithm, and proposed an improved swarm particle walking path optimization algorithm, and proved the effectiveness of the system through experiments. Literature [25] examined the application of artificial intelligence in intelligent decision support system to optimize English academic writing, and based on the literature review pointed out that artificial intelligence can improve grammar, sentence structure and make suggestions for clarity and articulation through tools such as Grammarly, which is beneficial for learners and learning institutions to improve the quality of English academic writing. Literature [26] introduces Wordtune, an AI writing assistant, which can understand the writer's ideas and make suggestions for optimization, which can effectively improve the efficiency of English writers, and by describing the role of Wordtune in English writing. Literature [27] created an intelligent English writing assistance system based on text mining technology, which utilizes advanced text mining techniques in order to analyze, evaluate, and enhance writing content, providing users with personalized guidance and support that promotes confident and effective communication in English. Literature [28] analyzed corpus-based English writing, while in order to improve the impact of the application of English teaching management system, the discussed system was designed with the goal of optimizing the teaching system and designing more practical functions with a view to enabling students to learn more, recording their English learning and enriching the content of the English course. Literature [29] investigated students' use of web-based digital writing assistants and examined the impact of different factors on students' adoption and use of digital writing assistants, indicating that the factors affecting students' behavioral willingness to use digital writing assistants are price value, effort expectation, etc. Literature [30] developed a writing aid scoring system for English learners, and by optimizing the system and designing an English writing aid evaluation scheme, it was revealed that the system has a high accuracy rate and good scoring results, which can be used for practical applications. The above

research proposes an intelligent system for English writing based on B/S model, machine learning, corpus, etc., which involves research in the fields of teaching, assessment, and academics of English writing, and on the whole, it confirms the validity of the proposed method, and it has a certain role to promote for both business English writing.

In this paper, in order to solve the recognition and efficient retrieval of image information in business English writing, a similarity calculation method based on word context vectors is introduced to improve the matching and retrieval of vocabulary recognition in images. After that, it is proposed to combine three-channel convolutional neural network (CNN) and bidirectional long and short-term memory network (BLSTM), embedded with the attention mechanism, to construct a new CBLSTM-Attention model network model, which can effectively deal with the fusion of image and text information. Then a personalized Pinyin-English thesaurus is built for different users by word alignment technology, the interactive recommendation of words is realized by using the language model, and a framework of an intelligent assistive system for English business writing is designed. Finally, the performance of the algorithm and the application effect of the model in this paper are analyzed.

2. Intelligent assistance system based on image recognition algorithm and language modeling

2.1. CBLSTM-Attention model structure

2.1.1. Word Vector Representation:

1) Data Preprocessing. There is a natural space character between words in English text, but Chinese text is different, so it is necessary to perform a lexical operation on Chinese text to separate words from words in Chinese text comments to form words with reasonable semantics. In this paper, the word separation tool used is the Chinese word separation tool written in Python language: jieba (jieba) word separation.

2) Text Quantization Representation. Text vectorization means that the result of word separation is represented by word vectors, and the word vector representation method used in this paper can transform text modeling into a machine learning problem, train large-scale corpus data, and output feature vectors with specified dimensions.

The CBOW model uses the context of current word W_t to predict the probability of current word W_t , and every word that appears in the context has the same influence weight on the occurrence probability of the current word; the Skip-gram model [31] uses current word W_t to predict its context probability. In this paper, we use Skip-gram model with a context window size of 5 to train word vectors on a large amount of unlabeled corpus, and the trained word vectors are used as model word vectors.

The objective function equation of the Skip-gram model is:

$$L = \frac{1}{T} \sum_{t=1-n}^T \sum_{i=j}^T \log P(W_{t+j} | W_t) \quad (1)$$

t is a parameter for the size of the training window, W_t is the current word, and W_{t+j} is the word appearing in the context. The word vectors generated by training represent the semantic information between words in a sentence.

2.1.2. Convolutional Neural Networks. Convolutional neural network is a type of feed-forward neural network, which has been widely used in the fields of image recognition, recommender systems

and nature processing. A typical network structure consists of an input layer, a convolutional layer, a pooling layer, a fully connected layer and an output layer.

The input layer is a vector representation of the input data, and the input layer matrix is denoted as:

$$e \in R^{n \times k} \quad (2)$$

where k is the dimension of the word vector and n denotes the length of the words that make up the text sentence. The convolutional layer uses different convolutional kernels to perform convolutional operations on the input matrix to extract the local features of the input and the result is denoted as:

$$c_i = f(W \cdot X_{i:i+h-1} + b) \quad (3)$$

where c_i denotes the i nd feature value after convolution operation; w denotes the weight matrix in the selected filter, where $W \in R^{h \times k}$, $h \times k$ denotes the selected filter size; b denotes the bias; f denotes the activation function. In this paper, the relu function is used as the activation function; $X_{i:i+h-1}$ denotes the length from the i th word to the $i+h-1$ th word in the text sentence. After the matrix $\{X_{1:h}, X_{2:h+1}, \dots, X_{n-h+1:n}\}$ of the input text is subjected to convolution operation, the convolution kernel feature vector map is obtained as:

$$C = [c_1, c_2, \dots, c_{n-h+1}] \quad (4)$$

Among them, $c \in R^{n-h+1}$. Since the text features are associated with the information before and after the word, in order to extract more features, this paper adopts a three-channel approach, i.e., three filters are selected for feature extraction.

Pooling layer is an important network layer of convolutional neural network, for the feature vector map obtained from the convolutional layer, the feature vector map is sampled through the pooling layer to extract the important feature information, this paper adopts the maximum pooling method for sampling, and the features obtained after the pooling layer are represented as:

$$c = \max(c_1, c_2, \dots, c_{n-h+1}) \quad (5)$$

2.1.3. Bidirectional long and short-term memory neural networks. The Long Short-Term Memory Neural Network (LSTM) is an improvement to the traditional Recurrent Neural Network (RNN) model, which solves the long-term dependence problem of RNN and the gradient explosion problem caused by too long sequences. The BLSTM model used in this paper [32-33], the bidirectional long short-term memory model at t moment is calculated as follows:

$$y_t = g(Vh_t + V's_t) \quad (6)$$

$$h_t = f(Ux_t + Wh_{t-1}) \quad (7)$$

$$s_t = f(U'x_t + W's_{t+1}) \quad (8)$$

where $x_1, x_2, x_3, \dots, x_i$ denotes the semantic expression of the input data, $h_1, h_2, h_3, \dots, h_i$ denotes the hidden layer set in the forward LSTM network, $s_1, s_2, s_3, \dots, s_i$ denotes the hidden layer set in the backward LSTM network, and $y_1, y_2, y_3, \dots, y_i$ denotes the semantic information after merging the forward and backward. For forward computation, h_t of the hidden layer is related to h_{t-1} ; for backward computation, s_t of the hidden layer is related to s_{t+1} .

2.1.4. Attention mechanisms. The attention model calculates a probability distribution of attention allocation, given the probability associated with any word in input sentence X , resulting in a probability distribution.

The attention value for each word within the text is calculated by the attention mechanism with the following formula:

$$u_t^i = v^t \tanh(W_1 h^i + b_w) \quad (9)$$

$$a_t^i = \text{soft max}(U_t^i) \quad (10)$$

$$c^t = \sum a_t^i h^i \quad (11)$$

where h^i is the text feature vector obtained by combining the three-channel CNN and BLSTM, W_1 is the parameter vector, b_w is the bias, and u_t^i is the hidden layer representation of h^i obtained by the neural network. a_t^i is the weight matrix obtained by normalizing the softmax function on u_t^i .

2.1.5. CBLSTM-Attention Model Construction. In this paper, we propose to combine three-channel convolutional neural network (CNN) and bidirectional long and short-term memory network (BLSTM), and combine the fused features with the attention mechanism to construct a new CBLSTM-Attention model network model. In this paper, the model first preprocesses the original corpus C and trains the word vectors through Word2Vec to construct the corresponding sentence vectors for each sentence; the constructed sentence vectors are inputted into the convolutional layer of the CNN and the BLSTM, respectively, and in order to achieve better experimental results, the convolutional windows are used in a combination of 3, 4, and 5, and the activation function used in the convolutional layer is RELU, with the step size set to 1. The sentence local features are extracted through the convolutional operation to extract sentence local features; then the features output from the three pooling operations are spliced together as part of the Attention input; at the same time, the constructed sentence vectors are used as the input to the BLSTM, and the two sets of features are obtained through the hidden layer LSTM units; finally, the outputs of the CNN and BLSTM are fused to be used as inputs to the Attention.

2.2. Efficient Search Algorithm Design

Adopting intelligent algorithms to improve the retrieval efficiency of specialized words or phrases can better assist the learning of students majoring in English business writing. This system sets up two types of algorithms: (i) to improve the retrieval efficiency of words or statements, a matching method based on similarity calculation is proposed; (ii) for the precious pictures of English texts existing in the field of chemical industry, a retrieval method based on deep learning is proposed.

2.2.1. Similarity-based word matching. In the field of English business writing, to better assist students in professional vocabulary searching, it is necessary to derive multiple words from a single word. In retrieval, this paper introduces a similarity calculation method based on word context vectors. The specific steps are divided into two steps.

1) Word vector generation. In this step, taking a keyword as the center, the context words of the keyword in the training set are counted to get the context vector of the keyword, and finally the vector matrix $TCV[n][n]$ is obtained.

2) Similarity Calculation. On the basis of the above vector matrix, the similarity of two words is calculated. The specific steps are: for any given two words, extract the corresponding context vectors from the vector matrix, and then calculate the cosine coefficient of the two vectors, i.e. the similarity value. The specific calculation formula is:

$$Sim[i][j] = \frac{\sum_{k=1}^n TCV[i][k] * TCV[j][K]}{\sqrt{\sum_{k=1}^n (TCV[i][k])^2} * \sqrt{\sum_{k=1}^n (TCV[j][k])^2}} \quad (12)$$

where Sim denotes the lower triangular matrix, $i \geq j$; $Sim[i][j]$ denotes the similarity of words t_i and t_j , $TCV[i], TCV[j]$ denotes the context vectors of words t_i and t_j , respectively, and $TCV[i][k], TCV[j][k]$ is the k th dimensional weight of the word context vector.

2.2.2. Image Retrieval Recognition. In order to further improve the role of this assistive learning system in vocabulary learning and better improve the application of English business writing, a fast retrieval algorithm based on the text of manuscript pictures is proposed, which enables users to retrieve the required pictures of related chemical materials through the pictures. In this paper, the combination of CBLSTM-Attention model is used to recognize the English pictures of manuscripts.

2.3. Interactive word-based recommendation models

2.3.1. Input processing. The spelling of Hanyu Pinyin at the semantic level lies in the combination of its syllables, and the following is the process of syllable input.

1) Syllable Slicing With the speed of input method technology, the existing phonetic and character conversion technology has been improved by leaps and bounds on the basis of rules and a large number of statistical techniques, and is becoming more and more humanized. Users in today's era have already developed the habit of frequently inputting simplified pinyin, half-full pinyin, or the whole sentence in the input process. In this paper, it is assumed that all the pinyin strings are correct, i.e., the correctness rate of the pinyin strings entered by the user's keyboard is considered to be 100%.

This system focuses on each pinyin syllable input in the interaction process, we adopt the method of cutting and recommending while inputting, and adopt the dynamic planning algorithm for the processing of simplified pinyin to achieve local optimization.

2) Feature Value Selection. Pinyin-word co-occurrence frequency: The system focuses on every Pinyin syllable input during the interaction process. When a Pinyin letter is input for the first time, the system is able to complete the currently input letter into a full spelling form by default. The model can use the frequency information of the co-occurrence of pinyin and Chinese characters to complete the pinyin syllables of the current single letter. The forward maximal matching method is used to match the read pinyin letters with the current statistical pinyin character co-occurrence list from front to back.

2.3.2. Establishment of the lexicon. Word alignment is used to build a thesaurus. The user's previously accessed text documents or domain documents related to the user's writing are stored locally and a personalized thesaurus is constructed with them. When the user inputs an operation, the system will match the user's input with the personalized thesaurus and give the user the output of the corresponding text.

First of all, English-Chinese and Chinese-English word alignment is performed to get the alignment information from all Chinese words to all English words or collocations and from all English words to all Chinese words or collocations, and then the word alignment information is merged with the Chinese-English and English-Chinese alignment information. The Chinese-English word alignment results are labeled with hanyu pinyin to build a phonetic-English thesaurus, and the word frequency information is retained.

2.3.3. Filtering and Sorting Relevant Cue Words:

1) n-Gram Language Modeling. A language model (LM) evaluates the plausibility of a particular word sequence in a language, giving the prior probability of a particular sentence in a given language.

The language model used in this paper is the most widely used n-Gram language model, which is called uni-Gram, bi-Gram, and tri-Gram when n is taken the values of 1, 2, and 3, respectively. For the pinyin syllable sequence $S = s_1, s_2, s_3, \dots, s_n$, the purpose of the phonetic-to-English conversion is to find the corresponding word sequence $W = w_1, w_2, w_3, \dots, w_n$ so that the generated word sequence satisfies:

$$w^* = \arg \max_w P(W | S, C) \quad (13)$$

where C is the contextual information of the user's writing process, i.e., the historical word sequence $C = c_1, c_2, c_3, \dots, c_n$. Since the candidate set is obtained according to the dictionary with fixed correspondences, the maximum probability of the current word under the condition of the existing word sequence can be calculated directly, i.e., the optimal candidate word can be obtained by the method of sorting the English language model, or according to the word frequency of the candidate word when there is no historical word sequence in the current input. For the n-tuple language model, the word sequence W is computed as follows:

$$P(W) = \prod_{i=1}^m p(w_i | w_1, \dots, w_{i-1}) \quad (14)$$

Language models are highly dependent on the size limit of the corpus, especially the n-Gram language model will have a stronger language description ability with the increase of the value of n, but at the same time there is a problem of data sparsity.

2) Lexical language model. In the input process of English, the current input word and history text combination is not necessarily a complete sentence and cannot form a complete syntactic tree, while given a word sequence, we can determine the lexical sequence it corresponds to. Considering the contextual information, the lexicality of the current word can influence the selection probability of the current word. Therefore, this paper introduces a lexical annotation model to help select the optimal word sequence. For lexical sequence W, its corresponding lexical labeling sequence T is chosen as:

$$T^* = \arg \max P(T | W) \quad (15)$$

Determine the candidate words in the current specific context. Annotate the lexical properties of the sentences composed of the current historical writing content and the candidate words with the probability of the lexical sequence T:

$$P(T) = \prod_{i=1}^m p(t_i | t_1, \dots, t_{i-1}) \quad (16)$$

We can get $p(t_i | t_1, \dots, t_{i-1})$ from the labeled training text, which is computed similarly to the language model $p(w_i | w_1, \dots, w_{i-1})$.

3) Dependency Syntax Modeling. Dependency analysis is a natural and language processing technique at the utterance level, which can not only obtain the syntactic structure of the sentence, but also the mutual dependencies between the individual words of the sentence.

For a sentence W that has been labeled with a lexical sequence T, the probability of the dependency tree D is defined as:

$$P(D | W, T) = \prod_i w_i \times f_i(D | W, T) \quad (17)$$

where $f_i(D|W,T)$ is the feature corresponding to each transfer state in the process of building the dependency tree.

4) Linear reordering method. In order to effectively utilize the various models proposed above, a linear reordering method is used to combine all the related words of each model and the sequence probabilities generated by the context and select the optimal related words from the top ranked candidates. The calculation is shown below:

$$P(R) = w_1 * P_w + w_2 * P_T + w_3 * P(D|W,T) \quad (18)$$

where $\sum_{i=1}^3 w_i = 1$ stands for the N-tuple language model, for the lexical language model, and for the dependent syntax model.

2.4. Overall framework design of intelligent auxiliary system

In order to improve the quality of English business writing, this paper develops a small English business writing assistance system. After users enter a few Chinese keywords into the search box of the system, the system will return English example sentences containing these keywords. For the sake of simplicity, the system is referred to as CEW system in the following paper, which firstly needs to pre-process the raw corpus data before entering into the database, and then performs the lexical processing after entering into the database. Users only need to search keywords directly in the search engine interface of the system. The main principle of the search engine is sorting and indexing, which greatly reduces the development difficulty and the problem of data storage. The system as a whole is divided into a three-layer design model, 1. User interface layer, using the form of a web page to input user queries and output related query results; 2. Chinese lexicon layer, the Chinese corpus and keyword lexicon, aligned with the English words; 3. Corpus layer, which is the core part of the system, the more the amount of corpus, the more accurate the results returned by the search. The work of CEW system is mainly divided into the following parts: building corpus, Chinese word segmentation, sorting, and searching.

1) Segmentation Process of Keywords and Corpus. A simple input box is set in the user interface, and after the phrase to be queried is inputted into this box, the CEW system first recognizes the phrase, and uses Jieba to separate it into Chinese words, so that the phrase is divided into some key words, and the accuracy of these key words being divided into words influences the system's ability to search for the first time, and the large-scale Chinese sentences in the corpus need to be separated into words, and the accuracy of Chinese words being separated into words influences the system's ability to search for the second time. The accuracy of Chinese word segmentation affects the system's ability to search for the second time, so Chinese word segmentation is especially important.

2) Chinese Segmentation Corpus Searching Process. The Whoosh search engine helps the CEW system to finish this work and get some sorted Chinese sentences with the highest similarity to the keywords.

3) Use of Chinese-English parallel corpus. The purpose of this system is to return the standard English sentences of mathematics, so it is necessary to output the English sentences corresponding to Chinese, which depends on the Chinese-English parallel corpus, but in this corpus, as the word meaning of English is variable, through the collection and collation of a large-scale English corpus and the introduction of WordNet related technology to improve the accuracy of the retrieved out of English, and to solve the problem of multiple meanings of the word and so on.

4) English Sentence Sorting Output Process. Through investigation, we found that there is an open-source search engine Whoosh developed by Python, which can achieve the effect of lightweight, easy

to install, greatly improving the efficiency of work, and has a good algorithm for English sentence sorting, and finally get the output results that satisfy the users.

CEW system work steps specifically see the process shown in Figure 1.

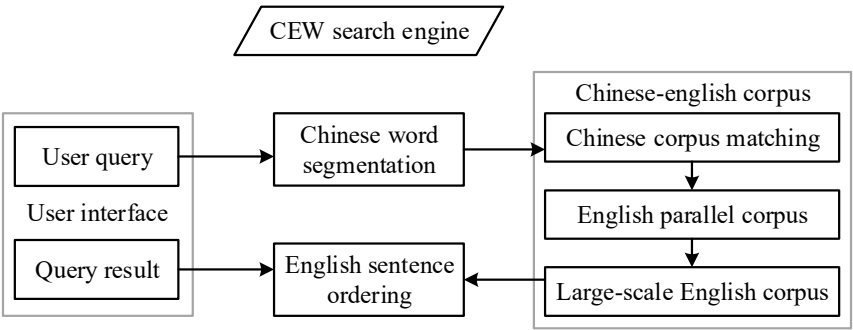


Fig. 1. CEW system workflow process

3. The Performance and Application Effect of English Business Writing Intelligent Assisting System

3.1. Comparison of Text Recognition Algorithms

In this experiment, a synthetic dataset is used as the dataset for the experiment, which contains a total of 632,172 images, of which 85% is used as the training set and 15% is used as the test set. Experiments are conducted on traditional text recognition methods, CNN text recognition methods, RNN text recognition methods and text proposed CBLSTM-Attention text recognition methods respectively. The training set of short text and long text are referred to as TS-1 and TS-2, and the test set of short text and long text are referred to as TS-3 and TS-4. The results of the experiments comparing the different text recognition algorithms and the visualization results are shown in Table 1 and Fig. 2, respectively. It can be seen that applying the recognition algorithm CBLSTM-Attention in this paper in the test set recognition accuracy, short text and long text than the traditional method to improve 25.56% and 29.82%, short text than CNN and RNN text recognition methods to improve 19.3% and 23.33%, long text to improve 16.53% and 28.96%, indicating that this paper's algorithm is able to effectively improve the accuracy of text recognition.

Table 1. Different text recognition algorithms compare experimental results

Data set		Text recognition accuracy(%)			
		Traditional model	CNN	RNN	CBLSTM-Attention
Training set	TS-1	71.32	80.95	79.8	98.58
	TS-2	64.3	75.83	65.83	96.24
Test set	TS-3	73.79	80.05	76.02	99.35
	TS-4	65.59	78.88	66.45	95.41

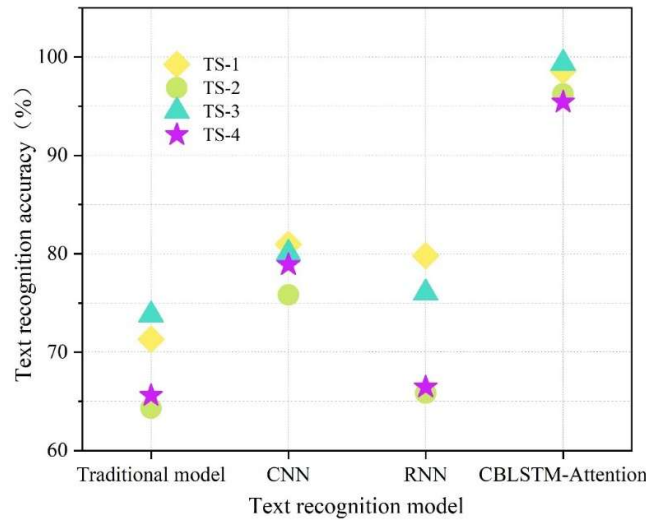


Fig. 2. Visual results of different text recognition algorithms

3.2. Comparison of Intelligent Assist System Performance Effect

Since word similarity is a highly subjective concept, there is no suitable test set for word similarity calculation. Therefore, we sort the set of synonyms of these 100 words relative to the query word by manual observation, set this sorting result as the reference sorting, and then compare it with the system sorting. Since the method of manual observation is too subjective, and at the same time the user is mainly concerned about whether the most similar word to the query word is ranked in the first place, we set that when the most similar word in the manually observed sorting result is ranked in the first place in the systematic sorting and there is at most one order inconsistency between the systematic sorting and the manual sorting, then we consider that the result of the systematic sorting is basically correct, and we record the manually determined most similar word of these 100 words position in the system sorting result. We use the following metrics to evaluate the algorithm:

1) The evaluation formula for accuracy is as follows:

$$precision = \frac{N_{correct}}{N_{total}} \times 100\% \quad (19)$$

where N_{total} is the total number of words tested, the systematic sorting result is considered to be correct when the most similar word observed manually is ranked first in the systematic sorting and at most one of the systematic and manual sorts is out of order. N_{total} is the number of system sorting results judged to be correct.

2) The reverse number of sorting (RR) is the inverse of the position of the correct result in the algorithm's return result, and the average reverse number of sorting is the average of the results of multiple calculations of RR, with the following formula:

$$MRR = \frac{1}{N} \times \sum_{i=1}^N \frac{1}{n_i} \quad (20)$$

where N denotes the total number of test words, and n_i denotes the manually determined most similar word for the i rd query word at position n_i in the system sorting result.

3) The number of exact matches of the query input in the example sentence retrieval system can directly reflect how common the input is in natural language processing. Since users mainly focus on the first word of the synonym recommendation, we record the number of exact matches of the first word of the synonym recommendation for each word in the test set in the ranking using the method

of this paper and the ranking using the baseline method, respectively, A and B. We tentatively refer to this method as evaluation method C.

3.2.1. Synonym Recommendation Effect. In this paper method CBLSTM-Attention is compared with RNN, CNN and baseline methods. The comparison result of synonym recommendation effect is shown in Fig. 3. It can be seen that this paper's method can select the most similar words of the query word very well compared to the 3 comparative methods such as CNN and baseline, and at the same time, the accuracy coefficient is improved by 0.32-0.6463 compared to the 3 comparative methods. It can be seen by the MRR value that this paper's method can give a reasonable degree of similarity between the similar words and the query word compared to the baseline method ordering and also improves up to 0.3887 compared to the accuracy of the three methods. It can be seen that CBLSTM-Attention can give the retrieval system an accurate ordering of the similarity of the words to be expanded. Through the C evaluation method, we find that the winning percentage of this paper's method is 96.21%, while the winning percentage of baseline's method is only 44.17%, which indicates that the common degree of the most similar words to the target word in the synonym recommendation produced by this paper's method is much larger than that of baseline and other methods.

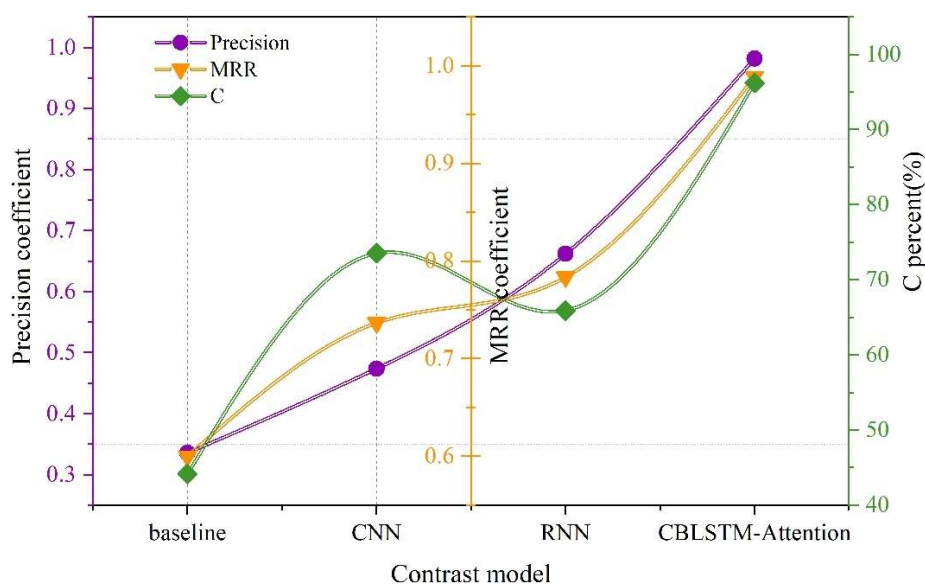


Fig. 3. Comparison results of synonyms recommendation

An example of similarity ranking in the context of ACL Anthology for the query words Outlook and Opportunity in wordnet and similarity in this paper is shown in Table 2. The synonyms of the word Outlook in the generic context are Opportunity, Future development, Growth potential and Market trend, and they all have a similarity of 1 under the similarity calculation of wordnet, which cannot distinguish the most similar word to the target word Outlook and obtain a similarity degree of ranking. The similarity of this paper gives better results here, the word Prospect is the most similar word to the word Outlook in this context, which is in line with the ACL Anthology usage under a large number of observations.

It is also observed that the word Ability does not appear in the ACL Anthology, which suggests that synonyms under the generic domain may not be at all similar to the target word in a particular contextual setting, and the two are not interchangeable. The case of word Prospect is the same as that of word Outlook, and the method of this paper can well distinguish the change of word similarity caused by the change of environment. Meanwhile, the method of this paper is also compared with the method without TFIDF, and it is found that the method without TFIDF has errors in the ordering of

some words with manual judgment, and most of the similarity values are lower than the method of this paper, so it is necessary to use TFIDF to calculate the weights.

Table 2. The similarity of outlook and opportunity in wordnet is similar to this article

Target word	Synonym candidate set	Wordnet similarity	baseline	CNN	RNN	CBLSTM-Attention
Outlook Prospect	Opportunity	1	0.7442	0.7096	0.7879	0.9936
	Future development	1	0.7126	0.7381	0.8223	0.9861
	Growth potential	1	0.4781	0.4414	0.5428	0.9781
	Market trend	1	0.5069	0.6415	0.6933	0.9664
	Potential	1	0.5251	0.5692	0.7823	0.9706
	Ability	1	0	0	0	0
	Expansion	1	0.4142	0.6058	0.4552	0.9322
	Innovation	1	0.2836	0.4126	0.3791	0.8921
	Sustainability	1	0.5192	0.7196	0.6789	0.9532

3.2.2. Word collocation effects:

1) Related experiments based on WordNet similarity calculation. Since word similarity is a subjective concept, our measure is still mainly a human judgment standard. The fitting degree of this paper's method with human subjective judgment is given by the formula:

$$FitRate = \frac{1}{N} \sum_{i=1}^N \left(1 - \frac{|SimRate_i - HumanRate_i|}{HumanRate_i} \right) \quad (21)$$

where $FitRate$ is the degree of fit between the manual judgment and the system calculated similarity, N is the selected experimental sample, $HumanRate_i$ and $SimRate_i$ are the manual judgment value and the system calculated value of the i th test sample.

In this paper, 8 word pairs are selected from WordNet for comparison, while the β value of the WordNet calculation method in this paper is set to 8, which is the value set after simple comparison. The manual evaluation value here is obtained by using the method of averaging multiple evaluations.

The experimental results of WordNet's similarity calculation are shown in Table 3. It can be seen that the degree of fit between the manual judgment value obtained by the calculation and the value calculated by the system is 95.92%, which is very well fitted with the manual judgment and can reflect the semantic similarity to a certain extent.

Table 3. The similarity of wordnet calculates the results of the experiment

Sample	Word 1	Word 2	CBLSTM-Attention	Human value
1	Economy	Technology	0.294	0.3006
2	Inflation	Innovation	0.3571	0.3911
3	Recession	Software	0.3125	0.3288
4	GDP	Hardware	0.3812	0.3663
5	Deflation	Block chain	0.4108	0.3905
6	Market	Automation	0.3231	0.3114
7	Budget	Virtual reality	0.3102	0.3028
8	Trade	Internet of things	0.3265	0.3292

2) Experiments related to phrase matching recommendation. The number of exact matches of a query input in an example sentence retrieval system can directly reflect the degree of commonness of that input in natural language processing. Therefore, we can automatically evaluate the commonness of

each collocation in collocation recommendation by the number of exact matches. Here Baseline is the average of the number of exact matches for a user query. We randomly selected 500 two-word-length user query from about 15,000 English query query that the system has been running for nearly three months. The experimental comparison results of the exact number of user query are shown in Table 4. It can be seen that although the average number of exact matches of all recommendations given by the system is smaller than the total average number of exact matches of the user query, the number of evaluated exact matches of the first recommendation given is much larger than that of the user's, which can be seen that the first recommendation given by the system is much more common than the user's in the ACL language environment, suggesting that the collocation recommendation given by the system may be a more desired collocation by the user. The first recommendation given by the system is much more common than the user's, indicating that the matching recommendation given by the system may be the one that the user needs more.

Table 4. Comparison results of user query accuracy

Recommended effects of CBLSTM-Attention	Average accuracy matching number
Baseline recommendation	452.87
First recommendation	578.46
All recommendations	259.15

3.2.3. Effectiveness of writing system applications. We invited 5 students majoring in Business English Writing as the subject of the experiment, and all students with different English proficiency levels were asked to write or revise their own English business articles using the English Business Writing Intelligent Assistance System designed in this paper. Each participant's essay was submitted to the same native English speaker for scoring, which included three criteria: "good", "medium" and "poor". In the experiment, the four students gave a total of 307 revisions to the paper. Table 5 shows the experimental results of 5 scholars using the English Business Writing Intelligent Assistance System to assist in writing or revising articles. After professional evaluation, 261 revisions were rated as "good" (85.02%), 31 revisions were rated as "medium" (10.10%), and 15 revisions were rated as "poor" (4.89%). From the experimental results, it can be seen that the acceptance rate of students' revision results with the help of the English Business Writing Intelligent Assistance System is more than 90%, and the system has substantially improved the quality of writing. The above experimental results and the actual situation of the use of the English Business Writing Intelligent Auxiliary System show the feasibility and effectiveness of the English Business Writing Intelligent Auxiliary English Writing System based on retrieval and recommendation.

Table 5. Use the CBLSTM-Attention system to assist in writing and revising articles

Student number	Articles	Experimental task	Modification number	Modified quality		
				Good	General	Worse
1	1	Writing and modification	43	38	4	1
2	1	Writing and modification	56	45	7	4
3	3	Writing and modification	78	61	12	5
4	4	Writing and modification	69	62	5	2
5	2	Writing and modification	61	55	3	3

3.3. Application Case Analysis of Intelligent Assisted Writing System

In order to verify the effect of CBLSTM-Attention applied to the teaching of English business writing, this paper selects 86 students from the 2024 grade of HZ University's Tiered Teaching A class as the experimental subjects for a 15-week teaching experiment, in which the experimental class and the control class are both 43 students. Before the experiment, this paper collected the English composition scores (total score of 100) of the entrance test of the experimental class and the control class in September 2024, and the results showed that the difference between the two classes was not obvious, which indicated that the study had some significance.

The experimental class was taught English composition using the CBLSTM-Attention assisted teaching model, while the control class still used the traditional teaching model. At the end of the experiment, this paper collected the final exam English composition scores of the experimental and control classes in January 2025. This study uses paired-samples t-test to compare the two scores of the experimental and control classes vertically; independent-samples t-test to compare the post-experiment scores of the experimental and control classes horizontally; and then analyzes whether there is a significant improvement in the English writing level of the experimental class compared with the control class with the help of SPSS software.

3.3.1. Comparison of performance on pre-tests. Before the beginning of the teaching experiment, this paper collected the English composition scores of two classes on the entrance test. The English composition scores of the September 2024 entrance test of the two classes are shown in Figure 4. It can be seen that the mean scores of the English composition scores of the entrance test of the two classes are very close to each other (69.8586 and 69.8593), which indicates that the students of the two classes had comparable writing levels before the experiment. The standard deviations of the two sets of data are also not very different, indicating that the internal dispersion of the composition scores of the two classes is also approximately the same. The results of the independent samples t-test of the English composition scores of the two classes on the September 2023 entrance test are shown in Table 6. The results show that the significance value is 0.8925, which is greater than 0.05. This shows that there is no significant difference between the English composition scores of the students of the two classes at the time of enrollment, and therefore the students of the two classes can be used as the object of comparative study.

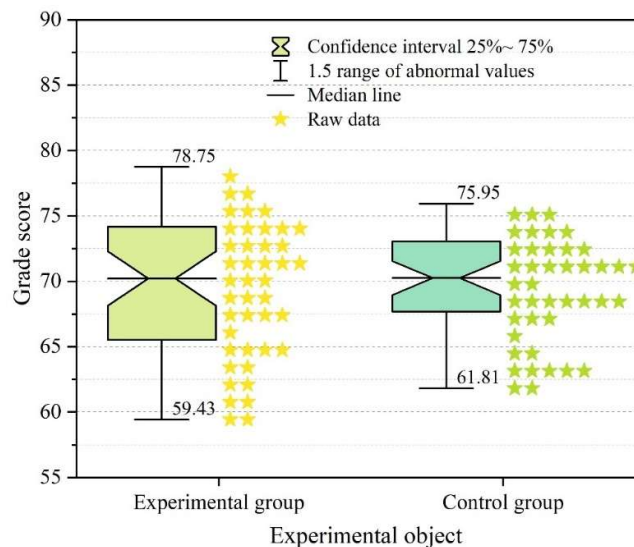


Fig. 4. The grades of English composition in the two classes in 2024

Table 6. Test results of the independent sample of the English composition

Test content	Index	Numerical value
Variance test	F value	0.0571
	Sig.	0.8925
	T	-0.0812
Mean test	Freedom	86
	Significance (double side)	0.9831
	Mean difference	-0.0007
	Standard error	0.1265

3.3.2. Comparison of post-test scores. After the teaching experiment began, the experimental class used the CBLSTM-Attention assisted teaching mode and the control class used the traditional teaching mode. The English composition scores of the two classes on the January 2025 final exam are shown in Figure 5. The composition scores of both classes improved, but the scores of the experimental class and the control class improved by 13.5902 and 3.6509 points from the entrance test scores, respectively, and the experimental class was 9.9393 points higher than the control class, and the difference between the scores of the two classes was very obvious, which shows that the experimental class improved its writing proficiency much more than the control class. The results of the

independent sample t-test for the final exam of January 2025 of the two classes are shown in Table 7. The significance value is 0.000, which is less than 0.05, and it can be seen that there is a significant difference between the composition scores of the final exam of January 2025 of the two classes.

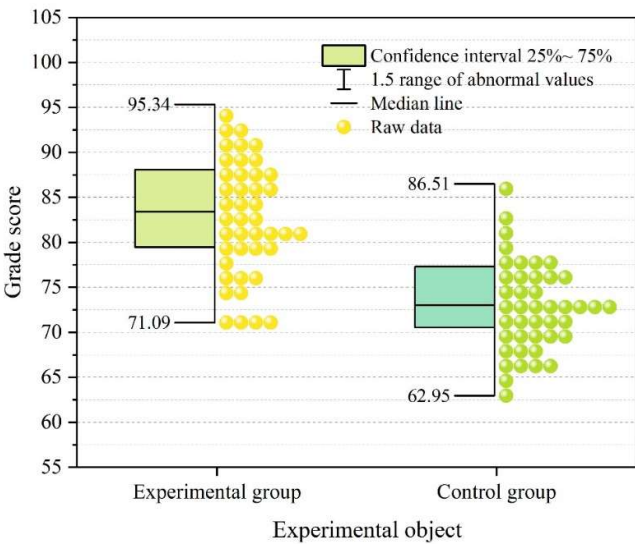


Fig. 5. After the test, the two classes of English composition results

Table 7. After the experiment, the independent sample T test of the test was tested

Test content	Index	Numerical value
Variance test	F value	27.3263
	Sig.	0.0000
	T	8.3614
Mean test	Freedom	86
	Significance (double side)	0.000
	Mean difference	9.9393
	Standard error	7.3751

3.3.3. Comparison of pre- and post-test scores in experimental classes. The results of pre- and post-test comparisons of the experimental class students' scores are shown in Fig. 6. The mean score of the experimental class' English composition score on the entrance test is 69.8586, and the mean score of the English composition on the final exam in January 2025 is 83.4488, and the mean score difference between the two scores is 13.5902 points. The results of the paired samples t-test for the two exam scores of the experimental class are shown in Table 8. It can be seen that the significant value is 0.000, which is less than 0.05, indicating that there is a significant difference between the composition scores of the two exams before and after, and that the intelligent writing assisted teaching has a significant impact on improving the writing level of the experimental class students in business English.

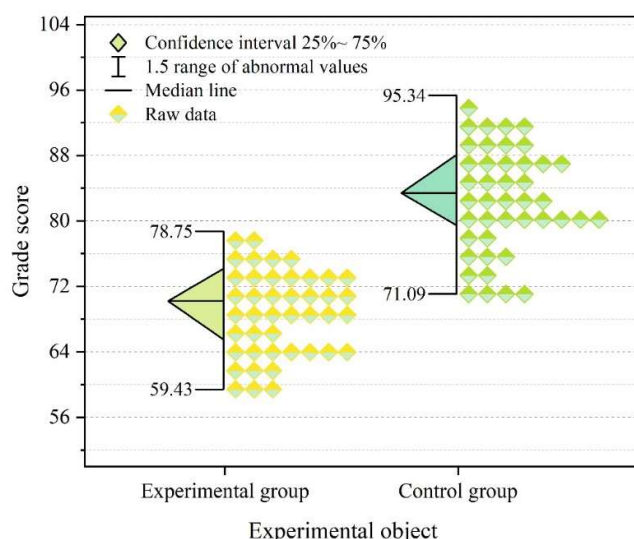


Fig. 6. The results of the test were compared with the students' scores

Table 8. Test results of the two test results of the experimental class

Test content	Index	Numerical value
Variance test	F value	23.3471
	Sig.	0.0000
	T	6.3906
Mean test	Freedom	86
	Significance (double side)	0.0000
	Mean difference	7.1542
	Standard error	6.2955

4. Conclusion

In this paper, we design an intelligent assistance system for English business writing that integrates image recognition and language modeling for the processing of data and language standardization needs in business English writing, and test the performance and application effect of the system.

The results show that: the algorithm in this paper can effectively improve the accuracy of text recognition, and its recognition accuracy is at least 18% higher than the three comparison algorithms. CBLSTM-Attention is significantly better than the other models in both synonym recommendation and word collocation, and the acceptability rate of the students' revision results for English business writing reaches more than 90%. Practical examples of the application of intelligent assistive writing system show that the quality of students' English business writing can be significantly improved.

This study verifies the application value of joint image-text modeling in English business writing intelligence and provides a new technical path for interactive writing assistance.

References

- [1] Wang, L., & Fan, J. (2020). Assessing Business English writing: The development and validation of a proficiency scale. *Assessing Writing*, 46, 100490.
- [2] Sun, B., & Fan, T. (2022). The effects of an AWE-aided assessment approach on business English writing performance and writing anxiety: A contextual consideration. *Studies in Educational Evaluation*, 72, 101123.

- [3] Sun, L., & Asmawi, A. (2023). The Effect of WeChat-Based Instruction on Chinese EFL Undergraduates' Business English Writing Performance. *International Journal of Instruction*, 16(1).
- [4] Chen, J. (2021). Research on the Effect of Peer Feedback Training in English Writing Teaching--A Case Study of Students in Business English Major. *English Language Teaching*, 14(6), 12-24.
- [5] Tsai, Y. R. (2021). Exploring the effects of corpus-based business English writing instruction on EFL learners' writing proficiency and perception. *Journal of Computing in Higher Education*, 33(2), 475-498.
- [6] Lin, C. J., Hwang, G. J., Fu, Q. K., & Chen, J. F. (2018). A flipped contextual game-based learning approach to enhancing EFL students' English business writing performance and reflective behaviors. *Journal of Educational Technology & Society*, 21(3), 117-131.
- [7] Feng, H., & Du-Babcock, B. (2016). "Business is Business": Constructing cultural identities in a persuasive writing task. *English for Specific Purposes*, 44, 30-42.
- [8] Tsai, S. C. (2019). Implementing interactive courseware into EFL business writing: computational assessment and learning satisfaction. *Interactive Learning Environments*, 27(1), 46-61.
- [9] Uba, S. Y., & Souidi, N. M. (2020). Students' Writing Difficulties in English for Business Classes in Dhofar University, Oman. *International Journal of Higher Education*, 9(3), 86-97.
- [10] Liu, M. (2019). Design of intelligent english writing self-evaluation auxiliary system. *Informatica*, 43(2).
- [11] Cui, L. (2018). Intelligent system for college English listening and writing training. *International Journal of Emerging Technologies in Learning (Online)*, 13(10), 121.
- [12] Tang, A., Li, K. K., Kwok, K. O., Cao, L., Luong, S., & Tam, W. (2024). The importance of transparency: Declaring the use of generative artificial intelligence (AI) in academic writing. *Journal of nursing scholarship*, 56(2), 314-318.
- [13] Pratama, R. M. D., & Hastuti, D. P. (2024). The use of artificial intelligence to improve EFL students' writing skill. *English Learning Innovation (englie)*, 5(1), 13-25.
- [14] Amanda, A., Sukma, E. M., Lubis, N., & Dewi, U. (2023). Quillbot as an AI-powered English writing assistant: An alternative for students to write English. *Jurnal Pendidikan Dan Sastra Inggris*, 3(2), 188-199.
- [15] Silaen, D., & Fitria, R. (2025). AI Writing Assistants: Insights from EMI Higher Education. *PROJECT (Professional Journal of English Education)*, 8(2), 306-316.
- [16] Athaluri, S. A., Manthena, S. V., Kesapragada, V. K. M., Yarlagadda, V., Dave, T., & Duddumpudi, R. T. S. (2023). Exploring the boundaries of reality: investigating the phenomenon of artificial intelligence hallucination in scientific writing through ChatGPT references. *Cureus*, 15(4).
- [17] Jacovina, M. E., & McNamara, D. S. (2017). *Intelligent Tutoring Systems for Literacy: Existing Technologies and Continuing Challenges*. Grantee Submission.
- [18] Liu, C., Hou, J., Tu, Y. F., Wang, Y., & Hwang, G. J. (2023). Incorporating a reflective thinking promoting mechanism into artificial intelligence-supported English writing environments. *Interactive Learning Environments*, 31(9), 5614-5632.
- [19] de Vicente-Yagüe-Jara, M. I., López-Martínez, O., Navarro-Navarro, V., & Cuéllar-Santiago, F. (2023). Writing, Creativity, and Artificial Intelligence: ChatGPT in the University Context. *Comunicar: Media Education Research Journal*, 31(77), 45-54.
- [20] Qin, F. (2022). College english intelligent writing score system based on big data analysis and deep learning algorithm. *Journal of Database Management (JDM)*, 33(5), 1-26.

- [21] Banawan, M., Butterfuss, R., Taylor, K. S., Christhlf, K., Hsu, C., O'Loughlin, C., ... & McNamara, D. S. (2023). The future of intelligent tutoring systems for writing. In *Digital Writing Technologies in Higher Education: Theory, Research, and Practice* (pp. 365-383). Cham: Springer International Publishing.
- [22] Wang, S., & Xu, H. (2018). Design of an Intelligent Support System for English Writing Based on Rule Matching and Probability Statistics. *International Journal of Emerging Technologies in Learning*, 13(11).
- [23] Fitria, T. N. (2023, March). Artificial intelligence (AI) technology in OpenAI ChatGPT application: A review of ChatGPT in writing English essay. In *ELT Forum: Journal of English Language Teaching* (Vol. 12, No. 1, pp. 44-58).
- [24] Guifang, L. (2021). Intelligent English writing system based on fusion of herding effect and artificial intelligence. *Journal of Intelligent & Fuzzy Systems*, 40(4), 6877-6887.
- [25] Marmoah, S., Adika, D., Haryati, S., & Yurni, Y. (2024). Leveraging AI to Optimize English Academic Writing (EAW) in Intelligent Decision Support Systems (IDSS). *Jurnal Ilmiah Universitas Batanghari Jambi*, 24(2), 1265-1270.
- [26] Zhao, X. (2023). Leveraging artificial intelligence (AI) technology for English writing: Introducing wordtune as a digital writing assistant for EFL writers. *RELC Journal*, 54(3), 890-894.
- [27] Fangzhou, Z. (2024, April). Design of an Intelligent English Writing Assistant System Based on Text Mining Technology. In *2024 IEEE 4th International Conference on Electronic Communications, Internet of Things and Big Data (ICEIB)* (pp. 363-367). IEEE.
- [28] Li, M., & Ye, W. (2024). Application of task processing system based on wireless sensor network in English writing intelligent assistance platform. *Measurement: Sensors*, 33, 101205.
- [29] Intiser, R., Nahid, M. H., Anwar, M. A., & Nahar, R. (2023). Adoption of AI-powered web-based English writing assistance software: An exploratory study. *AIUB Journal of Business and Economics*, 20(1), 90-101.
- [30] Lyu, J. (2022). Writing assistant scoring system for English second language learners based on machine learning. *Journal of Intelligent Systems*, 31(1), 271-288.
- [31] Enes Celik & Sevinc Ilhan Omurca. (2024). Skip-Gram and Transformer Model for Session-Based Recommendation. *Applied Sciences*, 14(14), 6353-6353.
- [32] EnverAkback & NedimMuzoğlu. (2024). An Efficient and Robust 3D Medical Image Classification Approach Based on 3D CNN, Time-Distributed 2D CNN-BLSTM Models, and mRMR Feature Selection. *Computational Intelligence*, 40(5), e70000-e70000.
- [33] Akback Enver. (2023). An efficient and robust supervised video hashing scheme based on a timedistributed CNN-BLSTM model and principal component analysis. *Multimedia Tools and Applications*, 83(21), 60965-60985.