

Contextual Analysis of Business English Translation Driven by Large Language Models

Xiaojing Li^{1*}

¹ Applied Foreign Language and International Education Department, Luohe Vocational Technology College, Luohe, Henan, 462000, China

Corresponding authors: (e-mail: mslijing1617@163.com).

Abstract With the acceleration of business globalization, accurate translation of business English has become an important guarantee for cross-cultural communication. This paper proposes a language model-based context analysis method for business English translation, which aims to improve the quality and translation speed of machine translation. The study by integrating CNN module and improved Attention mechanism, the model not only improves the translation quality, but also significantly accelerates the training speed. Experimental results in the WMT14 Chinese-English translation task show that the proposed model outperforms the traditional model in terms of BLEU score, up to 23.45, which is far more than the comparison model. Meanwhile, the model shows a significant advantage in training speed, taking only 6 hours and 10 minutes for a BLEU score of 17.58, with a convergence time about two-thirds shorter than that of the traditional model. The experiments also show that the integration of the model's attention mechanism and CNN module can effectively improve the translation quality and training efficiency, and the method can achieve higher translation quality in a shorter time. Taken together, this study can provide a more time-efficient and accurate solution in business English translation by optimizing the context analysis of machine translation.

Index Terms Neural Machine Translation, Language Modeling, Business English, Context Analysis, Attention Mechanism, Training Speed

I. Introduction

Over the years, China has been actively carrying out the “One Belt, One Road”, which, as a platform for global business exchanges and cooperation, has given language services a new historical mission. In terms of the social role of language, English, as an important medium of language communication, is widely used in international economic and trade activities, which has led to the formation of another specialized type of English translation “Business English” [1], [2]. Business English can be said to be a kind of special-purpose English, which mainly serves agents, producers, investors, consumers, etc., and is produced in the specific environment of related economic and trade communication activities, and used in the language of the specific environment [3]-[5]. With China's accession to the World Economic and Trade Organization (WETO) and the increasingly close ties with the world economy, the role of economic and trade English in strengthening the ties between China and the world has become more and more prominent, and the context is one of the key factors for the correct understanding of the language [6]-[8].

In economic and trade terms, the factors involved in context are mainly occasion, time, place law and other factors [9]. Usually, the context for the use of language mainly plays the role of “restriction and supplementation”, and the role played in translation is mainly in the interpretation of supplementation and screening restrictions [10], [11]. Context is an important condition for people to understand the speaker's semantics correctly, and it is also a key factor for realizing the purpose of communication. To correctly translate the speaker's meaning in the trade and to achieve the three principles of “faith, attainment, and elegance” in translation, the translator of business English must rely on the context [12]-[15]. Understanding the important role of contextual factors in business English translation is of great significance for mastering the correct way of translating business English.

The core of this study is to propose a novel method for contextual analysis of business English translation by combining the language model with the Attention mechanism in neural networks and Convolutional Neural Networks (CNN). It is investigated that by performing a fine-grained contextual analysis of the training data, the model is able to effectively utilize contextual information in the translation process and reduce the ambiguity of lexical translation. Further, the structure of the model is optimized by combining CNN and Attention mechanism, which makes the translation not only improved in quality, but also significantly shortened in training time. With this approach, the model is able to achieve higher translation quality in a shorter time, which is especially suitable for translation tasks

in the field of business English. In practical applications, this study uses the WMT14 Chinese-English translation task as an experimental validation, analyzes the performance of the proposed model in detail, and compares it with the existing NMT model. The experimental results show that the model proposed in this paper performs well in both BLEU score and training speed, and can effectively improve the efficiency and accuracy of business English translation.

II. Contextual Analysis Methods for Business English Translation Based on Language Modeling

II. A. Business English Translation Context Analysis Related Techniques

II. A. 1) Principles of machine translation

Language modeling is fundamental to the field of natural language processing. Let V be the word list, T be the length of the word list, and θ be the set of parameters, then for the sequence $Y = \{y_1, y_2, \dots, y_T\}$, $y_i \in V$, the language model can be expressed as:

$$P(Y; \theta) = P(y_1, y_2, \dots, y_T; \theta) \quad (1)$$

Machine translation is an important branch in the field of natural language processing, which also requires the use of language models for conversion. Machine translation usually uses an autoregressive language model [16], which decodes the source language words one by one when generating the target language, and takes the result of the last generation as the current input. Let $X = \{x_1, x_2, \dots, x_{T^X}\}$ is the source language sequence, T^X is the length of the source language sequence, $Y = \{y_1, y_2, \dots, y_{T^Y}\}$ is the target language sequence and T^Y is the target language sequence length:

$$P(Y | X; \theta) = \prod_{t=1}^{T^Y+1} P(y_t | y_{t-1}, X; \theta) \quad (2)$$

A maximum likelihood estimation approach is usually used to train machine translation models using a cross-entropy loss function.

II. A. 2) Encoder-Decoder Framework

Unlike common deep learning task processing (e.g., image processing, etc.), natural language processing employs an approach that mostly processes sequences. Currently commonly used models for machine translation use sequence-to-sequence [17] training methods, which can generate results with different lengths than the source language. Due to the general difference in the lengths of the source and target language sequences, this model structure is inherently suitable for the field of machine translation.

II. A. 3) Word embedding

As an important part of natural language processing, machine translation also needs to convert the language sequences used by human beings into multi-dimensional vector representations that can be computed by machines, which is also known as word embedding [18]. At the same time, word embedding is also the key to ensure that the model can successfully decode and generate accurate translations.

Traditional word embeddings use uniquely hot coding. This kind of sparse coding has the problems of “dimensionality disaster” and “semantic gap”. In order to solve these two problems, researchers have proposed word embedding models based on multidimensional vector representations, such as CBOW, SkipGram, Word2vec, GloVe, and so on.

Word embedding technology is widely used in the field of machine translation, and in practical use, word embedding vectors can be generated directly by using tools, or they can be generated by synchronous learning together with the model.

II. A. 4) Evaluation criteria for machine translation

A commonly used method for evaluating the effectiveness of machine translation is BLEU.

BLEU is a bilingual evaluation aid. It is mainly used to measure the accuracy of the target utterance, with larger values representing superior results.

Use l_c to represent the number of words in the candidate translation, and l_r to represent the number of words closest to l_c in the reference translation.

The calculation of BP is as follows:

$$BP = \begin{cases} 1 & lc > lr \\ \exp(1 - lr / lc) & lc \leq lr \end{cases} \quad (3)$$

BLEU4 is calculated as follows:

$$BLEU4 = BP \times \exp \left(\sum_{n=1}^4 W_n \times \log P_n \right) \quad (4)$$

where P_n is the proportion of human-translated translations to the total number of N-Grams in the machine-translated results, and $W_n = 1 / N$ is the maximum N-Gram order.

II. B.NMT in sentence-level topic contexts

II. B. 1) Attention-Based Machine Translation

In this section, an Attention [19] based NMT model (ANMT) is introduced to learn additional sentence-level topic context vectors for translation prediction. The recursive unit REC is first used to compute the hidden state proposal S'_i by using the previous decoder hidden state S_{i-1} and the previously fired target word $E[y_{i-1}]$:

$$S'_i = REC_1(S_{i-1}, E[y_{i-1}]) \quad (5)$$

where S_{i-1} is computed based on the original word context and the suggested topic context.

Second, the mechanism utilizes the hidden state suggestions S'_i to compute additional topic alignment weights ε_{in}^T for each T_n at time step i , which are weighted and summed with LTRs to obtain the topic context vector suggestions c_j^T :

$$\left\{ \begin{array}{l} \varepsilon_{in}^T = v_a^T \tanh(u_{S'_i}^{T'} + W_a^T T_n) \\ A_{in}^T = \frac{\exp(\varepsilon_{il}^T)}{\sum_{k=1}^N \exp(\varepsilon_{ik}^T)} \\ c_i^T = \sum_{n=1}^N A_{in}^T T_n \end{array} \right. \quad (6)$$

Intuitively, the purpose of NMT is to produce a sequence of target words with the same meaning as the source sentence, rather than to produce the same sequence of target topics. In other words, topic information can play a supplementary role in the translation prediction process. Therefore, $\lambda_i \in [0, 1]$ is computed, which is a gating scalar used to weight the expected importance of the source topic context for the next target word at time step i :

$$\lambda_i = \text{sigmoid}(w_\lambda s'_i + u_\lambda c_i^T) \quad (7)$$

w_λ and u_λ are model parameters. It is used sequentially at time step i to compute the topic context vector c_i^T , which is equal to λ_μ^T .

Then, using the topic context vector c_j^T as an additional input to a modified version of REC_2 , which is now based on the topic context vector c_i^T , one can compute the current decoder, using the original context vector c_i and the hidden state proposal s'_i , as follows hidden state s_i :

$$\begin{cases} z_i = \sigma(W_z c_i + W_z^T c_i^T + u_z s_i') \\ R_i = \sigma(W_r c_i + W_r^T c_i^T + u_r s_i') \\ s_i = \tanh(W_c c_i + W^T c_i^T + R_i \square (u s_i')) \\ s_i = (1 - z_i) \square s_i + z_i \square s_i' \end{cases} \quad (8)$$

Finally, the new decoder hidden state s_i is used to compute the probability of the next target word y_i while embedding the previously omitted target word $E[y-1]$ and two context vectors c_i and c_i^T :

$$\hat{P}(y_i | y < i, x) \propto \exp(L_d \tanh(L_y E[y_{i-1}] + L_w c_i + L_T c_i^T + L_s s_i)) \quad (9)$$

where L_o, L_y, L_w, L_T and L_s are projection matrices. Note that in addition to the extra subject context vector c_i^T , the original hidden state s_i needs to be replaced with the new decoder hidden state s_i .

II. B. 2) Themes in Transformer-based NMTs

In this section, it will be presented how the proposed CNN can be integrated into an existing NMT to jointly learn LTR and translation. The NM used in this study is the Transformer [20] based NMT.

First, the proposed CNN is used as an additional module of the encoder to learn LTR sequences from input source utterances. Second, an additional multi-head attention module is used to learn a new topic context representation of the target query based on the previous decoder layer. Compared to the original word-level context representation, the new topic context representation focuses on obtaining sentence-level topic information for translation prediction. Finally, the topic context vectors are used together with the original word context representation for predicting the target translation.

Formally, the CNN module of the encoder first learns the LTR sequence from the input source T statements. Then T is mapped to a set of key-value pairs $K, V = \{(K_1, V_1), (K_2, V_2), \dots, (K_M, V_M)\}$. In the decoder, multi-headed self-attention transforms the target queries of the previous decoder layers Q^i, K and V into H times:

$$\{Q_i^h, K^h, V^h\} = \{Q_i W_h^Q, K W_h^K, V W_h^V\} \quad (10)$$

Here $\{Q_i^h, K^h, V^h\}$ is the h -th subject header query, key and value vector. $\{W_h^Q, W_h^K, W_h^V\} \in \mathbb{R}^{d_{model} \times d_1}$ are model parameters. The thematic context of each thematic subspace can be represented by Eq. (11):

$$O_i^h = \text{softmax} \left(\frac{\{Q_i^h K^{h^T}\}}{\sqrt{d_k}} \right) V^h \quad (11)$$

Finally, the topic context vectors in the H subspace are concatenated into the current time step topic vector Q_i . According to the literature, it is known that both the new topic context vector Q_i and the original word context vector Q_i are used to compute the translation probability of the next target word via a linear, potentially multilayer function:

$$P(y_i | y_{<i}, x) \propto \exp(L_o \tanh(L_w O_i + L_T O_i)) \quad (12)$$

Here L_o, L_w and L_T represent the projection matrix.

II. B. 3) Explicit subject representation

Compared with traditional topic methods, the LTRs proposed in this paper speak implicit encoding of source topics via neural networks, which may be difficult to display the topic information encoded by these learned LTRs. In order to further understand the effectiveness of source topics, an explicit topic representation via Term Frequency Inverse Document Frequency (TF-IDF) is designed. Specifically, an input sentence of length J_g is considered as a document X_g and $TF-IDFTI_j$ is computed for each word X_g in X_j as shown in equation (13):

$$TI_j = \frac{n_{j_g}}{J_g} \times \log \frac{|G|}{1 + |g : x_i \in X_g|} \quad (13)$$

Here, $n_{j,g}$ denotes the number of occurrences of the j -th word in the input sentence d_g ; $|G|$ is the total number of sentences in the source language in the training data; and $|g : x_i \in X_g|$ is the number of source sentences in the training data that contain the word x_i . Then a fixed percentage of high TD-IDT words are selected and converted into word vectors to form a topic T sequence.

Compared with previous LTRs, T is explicitly extracted based on the TF-IDF method, and is thus called explicit topic representations (ETRs). Finally, the sequence of ETRs is utilized instead of the previous sequence of LTRs and integrated into the existing NMT architecture to enhance attentional alignment for business English translation prediction.

II. B. 4) Training models

The proposed topic-focused NMT is an integrated architecture that does not require any pre-trained topic information. In other words, it learns both LTRs and translations simultaneously rather than separately. Specifically, let θ denote the translation-related parameters and η denote the parameters used to learn LTRs. The loss function on a set of training examples $\{[x^k, y^k]\}_{g=1}^G$ is Equation (14):

$$T(\theta, \eta) = \sum_{g=1}^G \sum_{i=1}^{|y|} \log \hat{P}((y_i | y_{<i}, x; \theta, \eta)) \quad (14)$$

Therefore, the focus of training opposition is to find a group parameter (θ^*, η^*) to train the example $\{[x^n, y^k]\}_{g=1}^G$ in a set of training examples, which can be represented by equation (15):

$$T(\theta^*, \eta^*) = \arg \max_{\theta^*, \eta^*} T(\theta, \eta) \quad (15)$$

II. C. Support the implementation of context analysis translation models

We focus on the software implementation related to the key points of this topic and take it as the basis for further research and experiments. In the software, we will focus on how to implement language context analysis, context library management and its application in machine translation as the main breakthrough point of the language model in this study. Based on the basic word translation and sentence translation, and with context analysis as the basis, by optimizing the data structures and algorithms used in the software, a software prototype that can achieve high translation quality using context libraries will be developed. The translation of ordinary sentences is an implementation based on basic syntactic and grammatical analysis. As its algorithms are similar to those of other translation software, no more detailed explanations are prepared here.

The following is a brief description of the functional division, data structure design and algorithms of language models. The language model is divided into the main translation module, the extension module, and the intelligent learning and vocabulary maintenance module. The master translation module mainly performs general translation, that is, based on word meaning, grammar and syntax, to achieve translation that conforms to grammar and syntactic structure in the common sense. The extension module is mainly responsible for the implementation of contextual translation. Before the language model starts translation, the full text to be translated is retrieved and analyzed first, and an index library is established with the words and phrases of the full text as the index. It is convenient for the automatic acquisition of main keywords, the establishment and screening of contexts, and at the same time records the frequency of occurrence of keywords, compound words and phrases in each context and their offset positions (represented by bytes to determine the contribution of the character to the meaning extraction of words with multiple meanings). According to the content of the index, determine the genre and specialty to be used, and search for and initially determine the meaning and constraints of polysemy, colloquial expressions, slang, idioms, etc. in the context described in the article.

In order to apply the language context and cultural context corpora to the translation model, we have designed and optimized the involved language data structures many times. The definition of weight in the text is called "influence", which is the degree of necessity to take that meaning when the word appears, and it is related to the offset position of the word in the text. Computers can conveniently handle the calculation of offsets. Specifically, for each common word, it has different parts of speech, different forms of expression, and may also have different meanings. However, these meanings are not mechanically overlapped together, but are obtained based on its position (grammar) in the sentence, collocations, contextual conjunctions, keywords obtained during scanning, and comparisons with various weight levels.

To enhance the speed and improve the accuracy of translation, for the quantification and acquisition of some flexible conditions, an automatic analysis and question-and-answer approach is adopted. Users can set the genre,

subject category and affiliated category of the article in front of the language model to improve the accuracy and speed of meaning extraction. Users can "tune" while translating, that is, the method of learning while translating to add to the context corpus.

III. Attention-based machine translation model performance experiments

III. A. Experimental setup

In order to compare the quality of Attention-based NMT in neural machine translation modeling, this paper conducts experiments on the WMT14 Chinese-English translation task. The training dataset for the experiments is provided by the WMT14 Chinese-English translation task, which contains a total of 4.2 million parallel English- Chinese utterance pairs after preprocessing the corpus (deleting excessively long or short utterances, removing blank line operations, etc.), which involves about 110M English words and 120M Chinese words, and the newstest2014 English- Chinese database is used as a testing The newstest2014 English- Chinese database was used as the test set, containing 2500 utterance pairs, and the newstest2013 English- Chinese database was used as the validation set, containing 3200 utterance pairs. The parallel corpus training, test and validation sets are then processed using byte pair coding and a fixed lexicon is extracted from the training dataset with a lexicon size of 31000. All the models in this paper are coded in the python language and deployed on the tensorflow platform for implementation. All the experiments in this paper are performed on a hardware device with one NVIDIA 1080Ti GPU.

III. B. Analysis of the results of the comparison experiment

Table 1 shows the results of the comparison between the BLEU values of the Attention-based NMT machine translation model and the comparison model on the WMT14 Chinese-English translation task. The quality of the Attention-based NMT machine translation model proposed in this paper is substantially improved compared with the base comparison model. The large-scale Attention-based NMT machine translation model achieves the highest BLEU value of 23.45, and the small-scale model achieves a BLEU value of 22.56 that also exceeds most of the comparison models, so it can be seen that the Attention-based NMT machine translation model can significantly improve the quality of the neural machine translation model.

Table 1: WMT14 Chinese-English translation results

Model number	Feature	BLEU
1	The fourth layer of long-term memory network + dropout+ local attention mechanism + posnk	20.11
2	Convolution neural network coder + statement level target	21.89
3	Shared memory mechanism + GPV	21.56
4	4 layers of long-term memory network + attention mechanism + cluster search	22.08
5	Non-self-returning transformer	20.65
6	Greedy decoding algorithm + prediction mechanism	21.86
7	Batch alignment mechanism + filling limit	22.78
8	Attention+NMTmass	22.56
9	Small attention+NMT	23.45

Table 2 shows the change of BLEU value with time during training for Attention-based NMT machine translation model and comparison models: RNNSearch, Open NMT. The training speed of the Attention-based NMT machine translation model is twice as fast as that of the comparison model; the BLEU value of this paper's model is 17.58 at a training time of 6 hours and 10 minutes, while the RNNSearch model and the OpenNMT model need 15 hours and 20 minutes and 23 hours and 48 minutes, respectively, to reach a BLEU value similar to that of this paper's model. Similarly, our model takes only one-third of the convergence time of the comparison model, saving two-thirds of the training time, and our model reaches a BLEU value of 19.21 in 8 hours and 55 minutes, which is close to convergence, while the RNNSearch model only achieves a BLEU value of 8.73 in the same amount of time, and the Open NMT model only reached a BLEU value of 12.05. If we compare the models to achieve equal quality, the RNNsearch model needs 23 hours and 48 minutes to reach a BLEU value of 19.62, and the Open NMT model needs 53 hours and 37 minutes to reach a BLEU value of 19.61. Meanwhile, our model can reach a better quality at the initial time, at 3 hours and 51 minutes, our model can reach a BLEU value of 11.11, the RNNSearch model only reaches a BLEU value of 4.82, and the Open NMT model only reaches a BLEU value of 3.63, and this characteristic makes the model application in the contextual analysis of business English translation get a great advantage, the model can learn enough competence in the shortest time.

Table 2: Changes of BLEU scores along with the training time

Time	This model	RNNSearch	OpenNMT
53h37m	22.75	22.12	19.61
23h48m	21.59	19.62	17.79
15h20m	20.48	17.64	14.90
8h55m	19.21	8.73	12.05
6h10m	17.58	7.23	11.77
5h24m	14.55	6.32	7.76
3h51m	11.11	4.82	3.63

The BLEU value statistics are shown in Fig. 1, the horizontal axis is the training time step, and the vertical axis is the BLEU score, the curve of the model proposed in this paper is very steep at the initial time of training, which indicates that the BLEU score value of the model is increasing dramatically in this period of time, and it is easy to see that the curve of the BLEU value of the model proposed in this paper tends to be flattened after the steepness, which indicates that it reaches the convergence very quickly, while the curve of the comparison model's curve is flattening after rising slowly, the model proposed in this paper can greatly improve the training speed and reach convergence very quickly, which saves the training time.

In summary, the improvement work in this paper enables the model to obtain a more expressive representation. To achieve this goal, we integrate CNN into NMT to speed up the training and improve the quality. The advantages of the translation model in this paper are twofold. First, it significantly speeds up the training of neural machine translation models and can save a lot of training time compared to conventional systems. Second, it improves the translation BLEU score over the original neural machine translation model, which both ensures the translation quality of the model and saves the convergence time of the model.

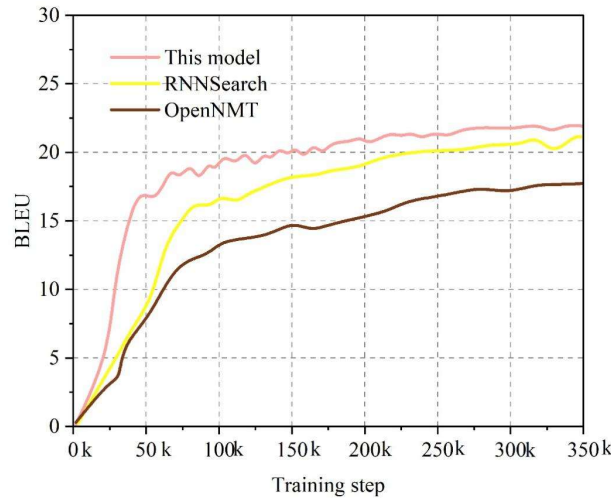


Figure 1: BLEU value statistics

III. C. Interpretability experiment analysis

In order to reflect the interpretability of the proposed business English neural network machine translation context analysis model, the attention distribution of sentences trained by the model is randomly selected in this experiment, and the attention distribution of the model is shown in Figure 2. The horizontal axis is the source statement "the same logic also applies to countries", and the vertical axis represents the target statement "the same logic also applies to countries". During the preprocessing, the sentence was segmented to obtain seven words of "same", "logical", "also", "applicable", "in" and "country", and a 7×7 matrix was formed with the target statement, and the matrix was grayscale processed to obtain the attention distribution map of the model. The higher the gray level of the graph, the stronger the attention information, and it is not difficult to see that the target words matched by each column are the darkest in color, that is, the attention guided graph convolutional network model has an alignment mechanism, which can reflect the interpretability of translation. In addition, the experiment also shows that the proposed model solves the problem of word order and translation selection. The word order problem can be seen from the positions of "same" and "same", and the model is trained to conform to the language habits of different

languages. In the translation selection problem, the grayscale column shows the correlation between different words, that is, the grayscale value indicates the degree of influence of its probability distribution value.

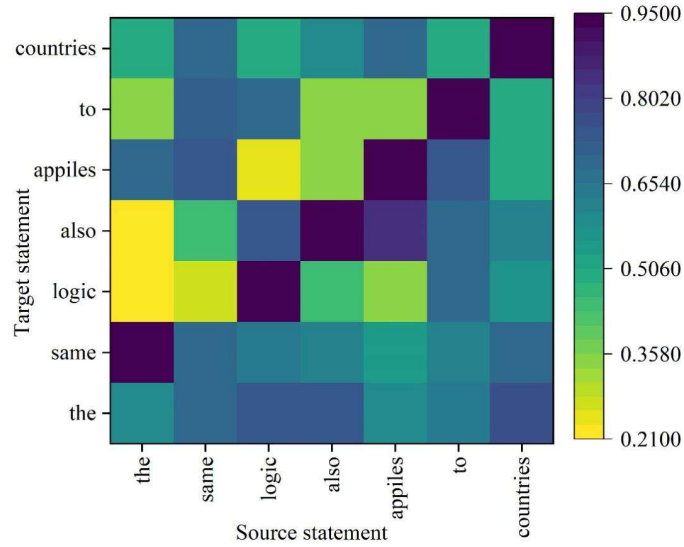


Figure 2: The attention distribution of the model

III. D. Analysis of ablation experiments

In the comparison experiments and interpretability experiments, it can be clearly seen that this paper's model has feasibility and superiority in machine translation compared with the traditional model, while the specific improvement methods based on this paper's model are not overly reflected in the experiments, therefore, this paper adopts a variety of types of ablation experiments, to verify that this paper's model has feasibility and superiority in the contextual analysis in business English translation.

In this experiment, three groups of processing strategies are firstly selected, which are: completely removing the attention mechanism, replacing the CNNs module with RNN, and removing the potential topic representation, and the comparison model adopts the complete model of this paper, and all the four models start to be trained in the same environment, and finally, the trained results are used to verify the effect of BLEU on machine translation, and the experimental results are shown in Table 3.

As can be seen from the results, the four sets of data differ greatly, indicating that the role of attention mechanism, CNNs module and potential topic representation in the NMT model is irreplaceable, and the complete model of this paper has obtained 23.56 and 24.12 BLEU in Newstest2013 and Newstest2014 datasets, respectively, which shows that the optimization of this paper is scientific which can provide a method for context analysis in business English translation based on language modeling.

Table 3: Experimental results

Model	Newstest2013	Newstest2014
Completely remove the attention mechanism	18.45	18.59
Replace the CNNS module with RNN	17.55	17.89
Remove potential themes	20.02	20.54
Complete model of this article	23.56	24.12

IV. Conclusion

Through the results of the comparison and ablation experiments, it can be seen that the language model-based context analysis method for business English translation proposed in this paper has achieved remarkable results in improving translation quality and speeding up training speed.

In the WMT14 Chinese-English translation task, the proposed Attention-based NMT model achieves a BLEU score of 23.45, which is far superior to other comparison models. In addition, the training time of the model is much shorter than that of the traditional model, achieving a BLEU score of 17.58 at 6 hours and 10 minutes, showing superior training efficiency and convergence speed.

From the results of the ablation experiments, it is clear that the role of the attention mechanism, the CNN module, and the latent topic representation in the model is irreplaceable. The experiments show that the performance of the model decreases significantly with the complete removal of the Attention mechanism, the replacement of the CNN module, or the removal of the latent topic representation, with BLEU scores of 18.45, 17.55, and 20.02, respectively, which are much lower than that of the full model, which is 23.56. This further verifies the effectiveness of the individual components in the model, especially the Attention mechanism plays a crucial role in the translation quality improvement plays a crucial role.

In summary, the business English translation context analysis method proposed in this paper not only improves the accuracy of machine translation, but also significantly improves the training efficiency of the translation model, providing a new solution for the field of business English translation.

Funding

This work was supported by Research on the Development Path of "Dual-element" Textbooks for Higher Vocational English in Accordance with the New Standards - Taking "Business English Communication and Customer Service" as an Example (No. WYJZW-2023IN0002), which is sponsored by the Teaching Steering Committee for Foreign Language Majors in Vocational Colleges of the Ministry of Education.

References

- [1] Gao, L. (2018). A Study on the Application of Functional Equivalence to Business English EC Translation. *Theory & Practice in Language Studies (TPLS)*, 8(7).
- [2] Tang, D. (2024). Study on the English-Chinese Translation of Business Contracts from the Perspective of Skopos Theory. *Journal of Theory and Practice in Linguistics*, 1(3), 1-9.
- [3] Lubić, A. A. (2022). Translation of the Passive Voice from English into B/C/S. *MAP Education and Humanities*, 2(2), 45-58.
- [4] Lu, W. (2021). Review and Reflection on Business English Translation Researches in China. *Education Research Frontier*, 11(2).
- [5] Türkmen, H. G. (2021). Systematizing the dialogue between translation studies and business studies: an interdisciplinary approach. *transLogos Translation Studies Journal*, 4(1), 81-99.
- [6] Hu, W. (2021). The application of cross-cultural awareness in business English translation. *International Journal of Social Sciences in Universities*, 4(3).
- [7] Al-Tarawneh, A., & Al-Badawi, M. (2025). Translation Quality and International Business Communication: Implications for Cross-Border Transactions and Investment Decisions. In *From Machine Learning to Artificial Intelligence: The Modern Machine Intelligence Approach for Financial and Economic Inclusion* (pp. 1089-1100). Cham: Springer Nature Switzerland.
- [8] Jia, Y., & Liu, L. (2023). The CE Translation Strategies Research of Cross-Border Ecommerce Products Under the Belt and Road Initiative. *Journal of Research in Social Science and Humanities*, 2(10), 33-36.
- [9] Ijoma, N. P. (2017). The Role of Context in Translation. *Journal of Modern European Languages And Literatures*, 8, 61-73.
- [10] Zhang, X. E. (2021). A study of cultural context in Chinese-English translation. *Region-Educational Research and Reviews*, 3(2), 11-14.
- [11] Laviosa, S., Pagano, A., Kemppanen, H., Ji, M., Laviosa, S., Pagano, A., ... & Ji, M. (2017). A contextual approach to translation equivalence. *Textual and contextual analysis in empirical translation studies*, 73-127.
- [12] Freitag, M., Foster, G., Grangier, D., Ratnakar, V., Tan, Q., & Macherey, W. (2021). Experts, errors, and context: A large-scale study of human evaluation for machine translation. *Transactions of the Association for Computational Linguistics*, 9, 1460-1474.
- [13] Liu, S., Chen, Y., Xu, K., & Lin, J. (2022). Emotional analysis of evaluation discourse in business English translation based on language big data mining of public health environment. *Frontiers in Public Health*, 10, 981182.
- [14] Yuan, Z., Jin, C., & Chen, Z. (2021). Research on language analysis of English translation system based on fuzzy algorithm. *Journal of Intelligent & Fuzzy Systems*, 40(4), 6039-6047.
- [15] Bi, S. (2020). Intelligent system for English translation using automated knowledge base. *Journal of Intelligent & Fuzzy Systems*, 39(4), 5057-5066.
- [16] Fanny Duce, Aurélie Névél & Karén Fort. (2024). "You'll be a nurse, my son!" Automatically assessing gender biases in autoregressive language models in French and Italian. *Language Resources and Evaluation*, 59(2), 1-29.
- [17] Jong Hyuk Lee & Min Young Kim. (2025). Manufacturing Quality Management Based on TimeGAN and Seq2Seq Models With Magnetic Press Machine Data. *International Journal of Control, Automation and Systems*, 23(4), 1199-1209.
- [18] Chanchal Patra, Debasis Giri, Sutanu Nandi, Ashok Kumar Das & Mohammed J.F. Alenazi. (2025). Phishing email detection using vector similarity search leveraging transformer-based word embedding. *Computers and Electrical Engineering*, 124(PB), 110403-110403.
- [19] Maryam Imani. (2025). Attention based multi-level and multi-scale convolutional network for PolSAR image classification. *Advances in Space Research*, 75(11), 7971-7986.
- [20] Arwa Alzughalbi, Adwan A. Alanazi, Mohammed Alshahrani, Ines Hilali Jaghdam & Abaker A. Hassaballa. (2025). Synergizing vision transformer with ensemble of deep learning model for accurate kidney stone detection using CT imaging. *Alexandria Engineering Journal*, 127, 357-373.