

Parallel Algorithm of Hierarchical Phrase Machine Translation Based on Distributed Network Memory

Guanghua Qiu, Henan University, China

ABSTRACT

Machine translation has developed rapidly. But there are some problems in machine translation, such as good reading, unable to reflect the mood and context, and even some languages machines cannot recognize. In order to improve the quality of translation, this paper uses the SSCI method to improve the quality of translation. It is found that the translation quality of hierarchical phrases is significantly improved after using the parallel algorithm of machine translation, which is about 9% higher than before, and the problem of context free grammar is also solved. The research also found that the use of parallel algorithm can effectively reduce the network memory occupation; the original 10-character content, after using the parallel algorithm, only need to occupy 8 characters, and the optimization reaches 20%. This means that the parallel algorithm of hierarchical phrase machine translation based on distributed network memory can play a very important role in machine translation.

KEYWORDS

Distributed Network Memory, Hierarchical Phrase, Machine Translation, Parallelization Algorithm

1. INTRODUCTION

In current Machine Translation, word order model is very important. There are many problems in translation models based on hierarchical phrases, such as cross syntax translation errors, word loss phenomena, and so on. (Tong & Zhu, 2016) is also very serious. Since the reordering effect of hierarchical phrase translation model mainly depends on the selection of hierarchical rules, and hierarchical rules are also the most essential feature of hierarchical phrase translation model different from other models, in this paper, we propose a joint selection method of hierarchical rules, which comprehensively utilizes various features generated in the process of translation, and solves the problem of correct selection of hierarchical rules to a certain extent, And it effectively solves part of the reordering problem based on hierarchical phrase translation model (Yang et al., 2017). Nowadays, the difficulty of statistical machine translation lies in the low syntactic and semantic components contained in the model (Bhadoria & Chaudhari, 2019). Therefore, when dealing with language pairs with large syntactic differences, such as Chinese to English translation, there will be a long distance reordering problem. Sometimes the result of translation can't be understood even though word to

DOI: 10.4018/IJISCM.2022010106

This article published as an Open Access Article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>) which permits unrestricted use, distribution, and production in any medium, provided the author of the original work and original publication source are properly credited.

word is correct. It can be said that the current mainstream statistical machine translation system is still at the bottom of the machine translation Pyramid (Li et al., 2017).

Machine translation refers to the automatic translation of text or speech from one language to another by computer, which belongs to the interdisciplinary research field of computer science and linguistics. In today's social life characterized by knowledge economy, the increasingly frequent international exchanges and the accelerating process of globalization make the total amount of cross language information exchange increase rapidly, and the natural language barrier between different countries and regions becomes more and more prominent. With the help of computer technology, people are eager to break down language barriers and realize smooth interpersonal and inter lingual communication. Therefore, machine translation research has great application value. Gradually improved machine translation technology has more and more extensive application in human social and economic life, and plays an extremely important role in the deeper and faster international information dissemination. Experts at home and abroad have also made a lot of research on hierarchical phrase machine translation (Liu et al., 2017).

Xiao Tong takes hierarchical phrase as the basic model and applies it to tree to string model reinforcement, which increases the space for translation derivation. In this framework, he studies the principles of machine translation and translation decoding technology (Buckley et al., 2019). In his experiment, the efficiency and performance of Chinese-English translation of machine translation have been greatly improved (Gulen et al., 2016) Jiangming introduces the shortcomings of machine translation at present. Aiming at these shortcomings, it uses Japanese case grammar and combines it with machine translation to mark different case information, which makes the language to be translated more obvious and easy to be analyzed. In addition, this method uses Japanese model, so it can achieve better results in Japanese translation, It proposes to use this framework in Japanese phrase translation Yegang hopes to integrate valuable syntax into machine translation, so that after translation, it can be more in line with the grammar and context of the country, and will not cause misunderstanding. Therefore, he puts forward three strategies to integrate bilingual into the machine translation, which integrates the phrases at the grammatical level in the way of hard constraints. The research finds that after the integration, the translated phrases and sentences are more perfect, and in complex long sentences, they are also more perfect It can play a certain role (Chen, 2019). These studies have certain reference value for this paper, but due to the lack of relevant research samples, there is no universal value, the research methods are not rigorous, and there are too few translation reference languages, so it can only improve certain theoretical reference.

In this paper, the hierarchical phrase machine translation model is described comprehensively, and the relevant theory of context free grammar is introduced. The training process of hierarchical phrase model, including rule extraction and rule scoring (Rui, 2017), is realized. The influence of the restriction of hierarchical phrase rule extraction on translation performance is verified by experiments. The decoder of hierarchical phrase model is also implemented, This paper introduces the data structure and efficient algorithm used in the decoder. Through the classification and analysis of hierarchical phrase types, the advantages of hierarchical phrase rules are obtained.

2. PARALLEL ALGORITHM OF TWO LEVEL PHRASE MACHINE TRANSLATION

2.1 Hierarchical Phrase Machine Translation

The process of translation is essentially a process of searching for the optimal solution. The translation system gives the translation with the highest translation probability from the huge candidate translation space. Modern statistical machine translation (SMT) was first proposed by IBM researchers. So far, it has developed from the original word based translation method to the phrase based translation method, and then to the syntax based translation method. Word based translation takes words as the unit of translation and ignores the word order relationship between words, resulting in poor effect of

translation order adjustment. The phrase based translation method extends the translation granularity from words to phrases, which can well solve the word order relationship within the phrase and improve the translation quality, but the performance of word order adjustment between phrases is poor (Feng, 2015). The syntax based translation model aims to overcome the defect that the phrase based translation method cannot use syntactic information, and takes the syntactic phrase as the basic translation unit.

Although the performance of the translation system has been improved in terms of word order adjustment, there are many problems with the syntax based translation methods. For example, the accuracy of syntactic analysis will have a great impact on the performance of the model (Lu et al., 2019); some non syntactic phrases beneficial to machine translation cannot be used in the translation process; when both the source language and the target language sentences are syntactic trees, the model is too complex. At present, even the best performance of the statistical machine translation system, there are still some problems, such as word mistranslation or missing translation, unreasonable syntax and improper long-distance word order adjustment. Statistical machine translation has entered a bottleneck period (Song & Huang, 2017).

The development of phrase based statistical translation methods has become mature. At present, many machine translation systems use phrase based translation models. When people use these translation systems for bilingual understanding, especially for long sentences and sentences with polysemy, the quality of translation is not satisfactory. Due to the use of non syntactic phrases as the basic translation unit in phrase based statistical machine translation, the model basically does not use any semantic information in the translation process (Sun et al., 2015). By introducing semantic knowledge into the translation model, semantic based statistical machine translation can solve deeper semantic problems in phrase translation, such as ambiguity and long-distance semantic constraints. Firstly, the verb object relation and subject predicate relation are extracted from the bilingual corpus. Secondly, in the case of two relations, the conditional probability based selection bias model and the topic based selection bias model are used for verb training (Zhang & Zhang, 2015). In addition, the verb is trained in the verb object model. Secondly, in view of the fact that the phrase based selection bias model is unable to estimate the selection bias for the relationship instances that do not appear in the training corpus, in addition to the simple unified value method, the word similarity method is also used to calculate the selection bias intensity of verbs for invisible words. Finally, the verb selection bias model is integrated into the decoder of the benchmark system as a new feature. The experimental results show that the statistical translation model based on verb selection bias can effectively improve the translation quality of verb object relationship instances and subject predicate relationship instances under long-distance semantic constraints (Biswas et al., 2018).

Due to the use of formal phrase as the basic unit in phrase based statistical machine translation (SBMT), most of the time the phrase does not have semantic information, so it can not identify the correct meaning of the polysemous words in the phrase, resulting in the translation errors of polysemy words. Word sense disambiguation is one of the effective ways to improve the accuracy of polysemy translation (Wan et al., 2015). In this paper, the coarse-grained level of word meaning, namely super meaning, is used to construct a word sense tagging device to label the super meaning of each word at the source end, so as to realize the high-level word sense disambiguation. Generally speaking, the word order of the source language and the target language will be very different. For example, Chinese and English, even if each phrase in the source language sentence is translated correctly, the position of the phrase is not in line with the word order of the target language, and the resulting translation is also incomprehensible. The existence of the reordering model is mainly to adjust the order of the generated translation, so that the word order of the translation can be in line with the word order of the target language to the greatest extent (Yang et al., 2019).

2.2 Distributed Network Memory

In order to solve the problem of limited communication and computing capacity of centralized network system, distributed network system is gradually developed. Usually, specific meanings are given according

to their respective research objects, such as aircrafts and robots in formation control, base stations in cellular communication systems, and individuals in human communities (Qiang et al., 2018). In the traditional centralized network system, each agent is usually required to obtain the information of all other agents, and then a centralized optimal control feedback is obtained by using the classical maximum principle (Namasudra & Roy, 2018). With the increase of demand and the enhancement of system functions, the number of agents is required to increase. In the complex geographical environment, the communication ability between agents is greatly limited. At the same time, the computing power of each agent is also greatly challenged. Therefore, the centralized network system cannot be used in many cases (Wu et al., 2015). Generally, the distributed network system does not need a centralized control unit, that is, the agent does not need to obtain all the information of other agents in the system, but a single agent in the system only needs to obtain the information of some agents in the system, then calculates, and achieves the goal through the cooperation of the agents to achieve the global control (Kitouni et al., 2018; Qi, 2018).

Distributed network system provides a feasible solution to solve the problem of communication and computation limitation in traditional networked control system, and it has the advantages of computational complexity, system limitation, scalability and robustness:

$$x(k+1) = Ix(k) + Jv(k), k = 1, 2, \dots \quad (1)$$

The quadratic performance index of the above problem is as follows:

$$K = \sum_{k=1}^{\infty} [x^i(k) J x(k) + r^i(k) c J] \quad (2)$$

where the weighted matrix Q is:

$$Q = \frac{1}{2a^2 r^{-1}} \left(\frac{2b^2}{a^2 r^{-1}} p - t \right)^{-1} \left[a^2 r^{-1} t^2 + 2(1-b^2)t \right] \quad (3)$$

If $a \in [-1, 0] \cup [0, 1]$, then the constant feedback gain matrix K of the system is:

$$K = \frac{a}{2br} t \quad (4)$$

The constant feedback gain matrix of the system is of the same type as the Laplace matrix, so the optimal control of the distributed network system with only the information of adjacent nodes can be obtained by introducing K into the formula. Available:

$$\lambda_x(ct_n - t) > 0 \quad (5)$$

Therefore:

$$Q = \frac{1}{2a^2 r^{-1}} \left(\frac{2b^2}{a^2 r^{-1}} t - L \right)^{-1} \left[a^2 r^{-1} L^2 + 2(1-a^2)L \right] \quad (6)$$

By introducing the parameter and the weighted matrix Q, we can get:

$$\left(\frac{2b^2}{a^2 r^{-1}} I_x - t \right) Q = \frac{1}{2} t^2 + \frac{1-b^2}{a^2 r^{-1}} t \quad (7)$$

It can be obtained by formula:

$$Q^2 + \frac{2(1+b^2)}{a^2 r^{-1}} Q + \frac{(1+b^2)^2}{(a^2 r^{-1})^2} I_x = \left(Q + t + \frac{1-b^2}{a^2 r^{-1}} I_x \right)^2 \quad (8)$$

Combined with the correlation formula, we can get:

$$K = \frac{a}{2rb} \left[Q + t + \frac{1-b^2}{a^2 r^{-1}} I_x - Q - \frac{1-b^2}{a^2 r^{-1}} I_x \right] = \frac{a}{2br} t \quad (9)$$

Therefore, the constant feedback gain matrix of the system is of the same type as the Laplace matrix, and the distributed network system with only the information of adjacent nodes can be obtained by introducing K into the formula.

2.3 Hierarchical Phrase Machine Translation Model

Word alignment describes the corresponding relationship between source language words and target language words. Some people take the lead in introducing the idea of word alignment into statistical machine translation. This method is favored by researchers, and the current statistical machine translation is based on it (Zhang, 2015).

For a simple sentence, when it is translated, it can be translated into the target language word by word. In this case, it can be seen that there is a word correspondence between the source language and the target language. Of course, there may also be words in the source language that have not been translated, and words in the source language corresponding to several words after translation (Wang et al., 2015). This phenomenon can be described in several aspects: the probability of word to word translation, the change of word position in the process of translation, the probability that a word is translated into multiple words, etc.

Suppose there is a source language sentence and a target language sentence, define an alignment set, ax means that the X source language word corresponds to the ax target language word. Each source language word corresponds to at most one target language word, and one target language word may correspond to multiple source language words (Sun et al., 2019). Such alignment can be called one to many one-way alignment from the target language to the source language. Therefore, for a certain result a , we can express it by formula:

$$P(K / T) = \frac{\varepsilon}{(y+1)^n} \prod_{i=1}^n \sum_{j=0}^y P(K_i | T_j) \quad (10)$$

The most likely alignment of K and t can be obtained by the formula:

$$a = \arg \max_a \prod_{i=1}^n P(K_i | T_j) \quad (11)$$

Given a phrase training set, according to the above two formulas, EM algorithm can be used to estimate the parameters of the model.

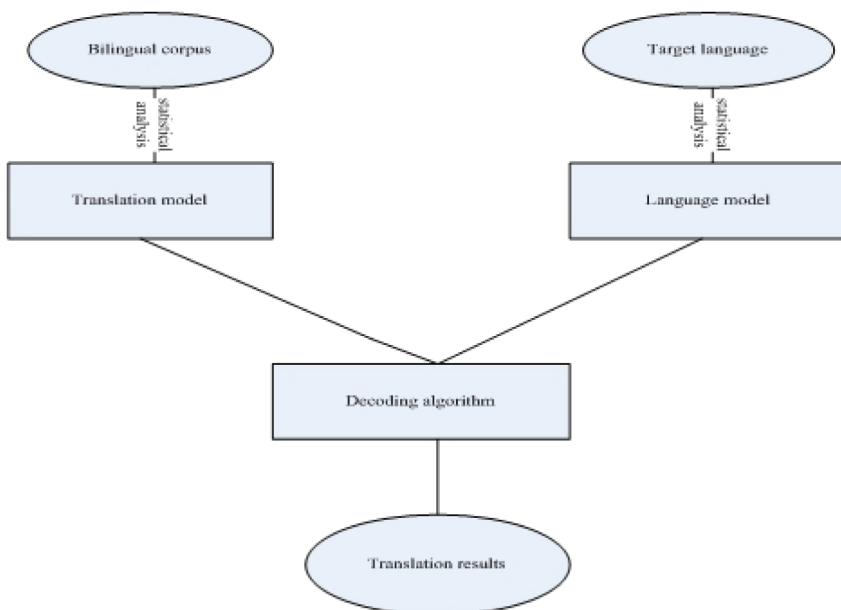
The translation model reflects the probability of translation from the source language to the target language, and the language model represents the rationality of the target language rules, and reflects the rationality and fluency of the target language (Hu et al., 2018). The decoding algorithm decodes the input source language according to the established model and training parameters, and finally gets the translation results. The quality of the algorithm directly affects the search space and time efficiency of decoding. The phrase based statistical machine translation method is superior to the word based statistical machine translation method in terms of performance because it can better grasp the local context dependence. Compared with syntactic based statistical machine translation, phrase based statistical machine translation has good generality and small search space (Deng & Yang, 2016). Machine translation system is generally divided into three parts, as shown in Figure 1.

2.4 Level Phrase Decoding

On the whole, the main function of the decoder is to search the translation model and translate the given sentence. Specifically, the decoder needs to do the following work.

In this way, the storage space can be saved and the time can be greatly improved. In decoding, we need to convert the string into a number for operation, and at the end of decoding, we need to convert the number into a string for output. Effectively represent the rules used in the decoding process. In decoding, a large number of rules need to be searched. Therefore, effectively representing these rules is conducive to saving memory and speeding up the search. We use prefix tree (tree) to store rules, that is, before decoding, we need to load all the results of previous fix tree (tree) that may be used into memory (Ma et al., 2018).

Figure 1. Machine translation system



In decoding, we need to quickly retrieve the rules needed to translate the current sentence. The reason why we use prefix tree for storage is that prefix tree can not only save memory, but also can quickly find the rules we need with the algorithm. The core algorithm of decoder is to use dynamic programming algorithm, which establishes a two-dimensional table, and then deduces, prunes and searches in the table, so as to finally translate the final result (Zhai et al., 2019).

The data structure and algorithms mentioned above are introduced in detail.

Digitization: in order to save storage space and improve retrieval speed, we digitize all the strings used in the translation process. From the implementation point of view, string class needs to construct objects and release objects, while integer numbers are built-in data types, so operating numbers are much faster than string operation. The specific method is to use two hash tables to store the mapping of string to number and number to string respectively. In decoding, we need to convert the string into a number for operation, and at the end of decoding, we need to convert the number into a string for output. In addition, we need to determine whether the character is a non terminal character in the program. Therefore, we map the terminator to a positive number and the non terminator to a negative number. In this way, we can judge whether the character is a terminator or not according to the positive and negative of the number.

CKY algorithm: CKY algorithm is a bottom-up search algorithm, which requires the grammar to conform to Chomsky normal form, that is to say, the right side of the production is either two non terminals or one terminator. Since any context free grammar can be converted into an equivalent grammar, CKY algorithm can be applied to any context free grammar.

Generally speaking, there is at least one grammar with linguistic significance in the source language and the target language, which is mainly obtained by parsing the source language and the target language by the syntactic analyzer. In order to use such syntactic information in statistical machine translation model, grammar is used in syntax based translation model, and the corresponding relationship between source language and target language is established by synchronous grammar rules. The purpose of using synchronous context free grammar in statistical machine translation is that it can maintain the integrity of phrases with syntactic structure.

3. PARALLEL ALGORITHM EXPERIMENT OF HIERARCHICAL PHRASE MACHINE TRANSLATION

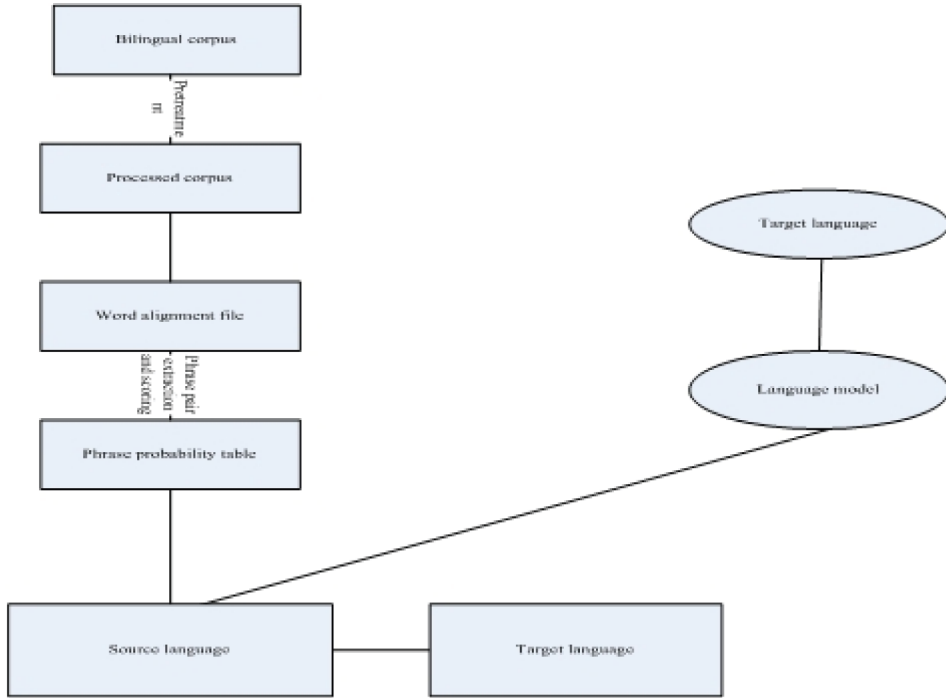
3.1 Structure of Machine Translation System

The machine translation system constructed in this paper is based on Moses open source software. The benchmark system consists of two modules: training and decoding. The training module uses Giza ++ for bidirectional training to obtain word alignment file, then extracts phrases and obtains phrase translation probability table, and uses training language model. The decoding module divides the input sentences into phrases, and then searches for the optimal combination according to the existing phrase translation probability table. Its structure is shown in Figure 2.

3.2 Data Preprocessing

Hierarchical phrases that need to be translated need to be transformed into a form acceptable to Giza ++. In this paper, according to the frequency of words, each word is assigned an index number to generate two word files. The file contains the word index number, the word, and the number of times the word appears. The index number corresponds to the empty word. After preprocessing, there will be three paragraphs in a sentence. The first paragraph is the original sentence, and the second and third paragraphs are expressed in two common languages, such as English and Chinese.

Figure 2. Translation structure



3.3 Level Phrase Translation Alignment

The word alignment results obtained by Giza ++ ignore many to many, so they can not reflect the relationship between source language and target language. Therefore, on the basis of word alignment, phrase alignment is carried out, which is conducive to better grasp the context dependence. Phrase extraction is based on the optimized bidirectional word alignment results. All rectangles with alignment points as vertices in the alignment matrix are extracted. The condition is that the target words aligned with the source language words in the row range of the rectangle are in the column range of the rectangle, and vice versa. When extracting phrases, phrase pairs must be compatible with word alignment:

$$p(x, y) = \frac{1}{n} * i, j \quad (12)$$

where I is the number of times x occurs in the sample and j is the number of times y occurs in the sample.

The model depends on the information of context. Suppose we give a characteristic function to restrict each feature: the expected probability value is equal to the empirical probability value, as follows:

$$p(f_a) = \sum_{x, y} p(x) p(y|x) f(x, y) \quad (13)$$

Our goal is to find the optimal $P(x, y)$. First, we get the input samples with fixed format required by the maximum entropy from the corpus, and then train the model by the maximum entropy training tool. Using the maximum entropy model, we can return the value of conditional probability according to the maximum entropy model.

3.4 Score of Parallel Algorithm for Hierarchical Phrase Translation

In this paper, the possibility of each phrase belonging to each topic in the phrase translation probability table is evaluated, and then the topic model is applied to the translation system. The topic score of a phrase is determined by its sentence. When extracting phrase pairs, the sentence number of each phrase pair should be recorded. Finally, the topic score of the phrase can be obtained by the topic score of these sentences. After getting the topic probability of all words, this paper calculates the score of each topic according to the topic probability of each word in the sentence. The scoring formula is as follows:

$$score(p_x, a) = \sum_{n=1}^M \left(t(p_x | r_n) - p(p_x | r) \right) \quad (14)$$

where p is the average of all word occurrences.

4. PARALLEL ALGORITHM ANALYSIS OF FOUR LEVEL PHRASE MACHINE TRANSLATION

4.1 Level Phrase Length

The length of hierarchical phrases has a certain impact on translation performance. Generally speaking, the initial length of a phrase is limited to 10. If the length of a phrase is too long, the decoding efficiency of the phrase will be reduced. If the phrase is too short, the expression ability of the phrase will be weakened. Therefore, we select the phrases with the length of 5 to 15 to conduct the experiment, and the experimental results are shown in Table 1.

From figures 3 and 4, we can see that the performance of translation is changing as the initial phrase grows. From the initial 5 to 10, the translation performance changes most intensely. When the number of phrases exceeds 10, the change tends to be gentle. This shows that when the initial phrase length exceeds 10, the decoding efficiency and translation performance are the most balanced.

4.2 Impact of Parallel Speed on Machine Translation Algorithm

Different algorithms have a great impact on the speed of machine translation. In order to synchronize the translation speed and performance of hierarchical phrases, we have made statistics on the translation practice of hierarchical phrases, as shown in Table 2.

From Figure 5 and Figure 6, we can see that the parallel algorithm has certain advantages over the previous general algorithm in hierarchical phrase translation, and it becomes more and more obvious with the change of the length of the hierarchical phrase. We can see from the figure that when the length of the hierarchical phrase is 10, the calculation speed of both sides is basically the same, but

Table 1. The influence of hierarchical phrase length

	5	7	10	12	15
General algorithm	25.9	26.5	26.9	27.1	27.3
Parallel algorithm	25.9	26.3	26.7	26.8	26.8

Figure 3. Translation level of different algorithms

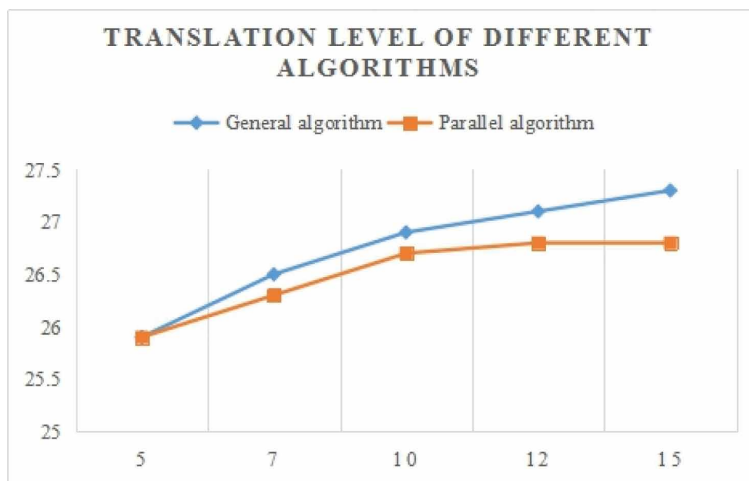


Figure 4. Phrase length translation

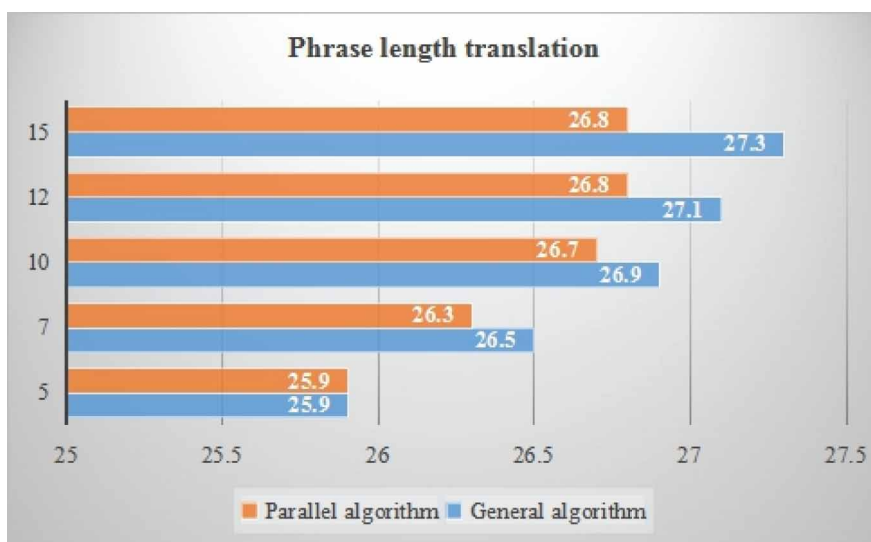


Table 2. Translation speed

	10	12	15	17	20
General algorithm	1.23	1.25	1.26	1.37	1.44
Parallel algorithm	1.17	1.19	1.23	1.25	1.26

Figure 5. The influence of translation methods on speed

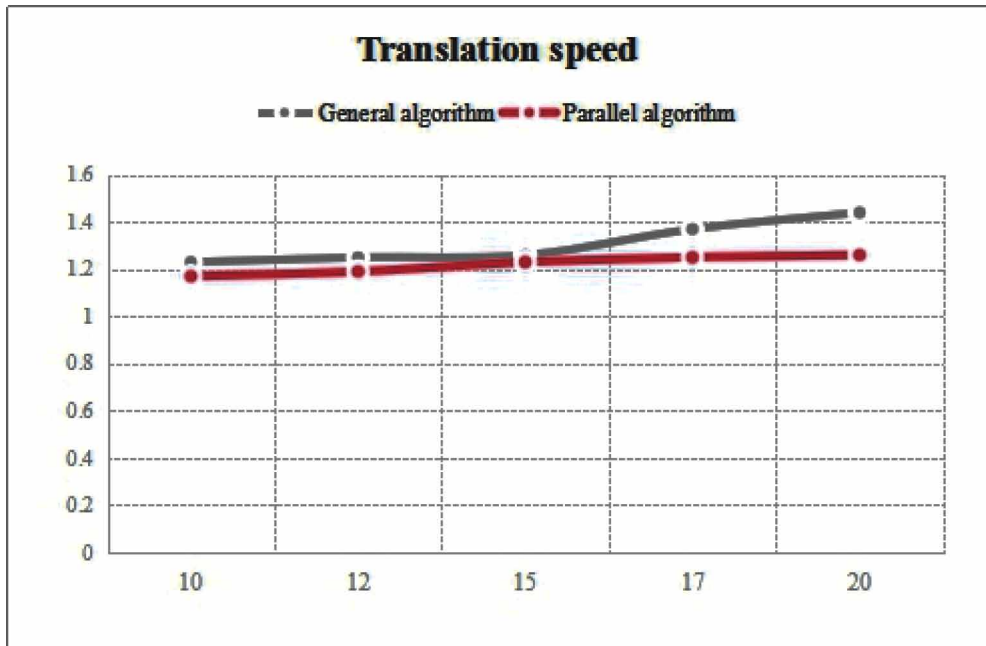
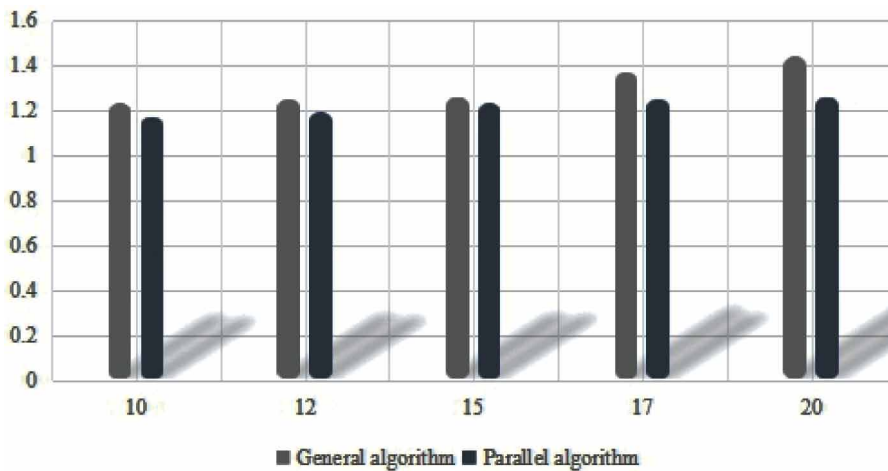


Figure 6. Translation speed difference



when the length of the hierarchical phrase reaches 20, the difference between the two sides reaches the maximum. When the length of the parallel phrase is too long, there is an advantage of the algorithm.

4.3 Influence of Cohesion and C Value on Phrase Translation

C value is an important part in phrase translation. In this paper, we reduce the probability table of phrase translation according to the C value of the source language without considering the length of the phrase. As shown in Table 3.

Table 3. Cohesion and Translation

C value	Phrase expression size	BLUE
0	100%	0.402
0.9	83%	0.418
1.5	72%	0.423
1.9	64%	0.416
2.5	57%	0.414
2.9	43%	0.403

It can be seen from Figure 7 and figure 8 that the blue evaluation can be up to 0.11 higher than the benchmark system, while the phrase translation probability table is only 69% of the original. Moreover, when the phrase translation probability table is reduced to 43%, the blue evaluation is still slightly improved compared with the benchmark system.

4.4 Application of Parallel Algorithm in Model

We grade the sentences in the test set on the topic. If the score of a certain topic is the highest, we classify the sentence as the topic. In this way, we can classify the test set into four categories with larger data sets for testing. After comparing the changes after the parallel algorithm, the specific data are shown in Table 4.

From Figure 9 and figure 10, we can find that when testing the four types, we take a threshold for the relevant data, and filter out the hierarchical phrases that are smaller than the threshold, which has a certain degree of optimization for machine translation methods. In Table 4, we can see the threshold conditions of each test; after filtering, different thresholds can improve the translation results. In the second and fourth groups, the improvement is higher than the other two groups. This is because the algorithm used in these two groups is more excellent, so the translation results are improved more greatly.

Figure 7. C value and Translation

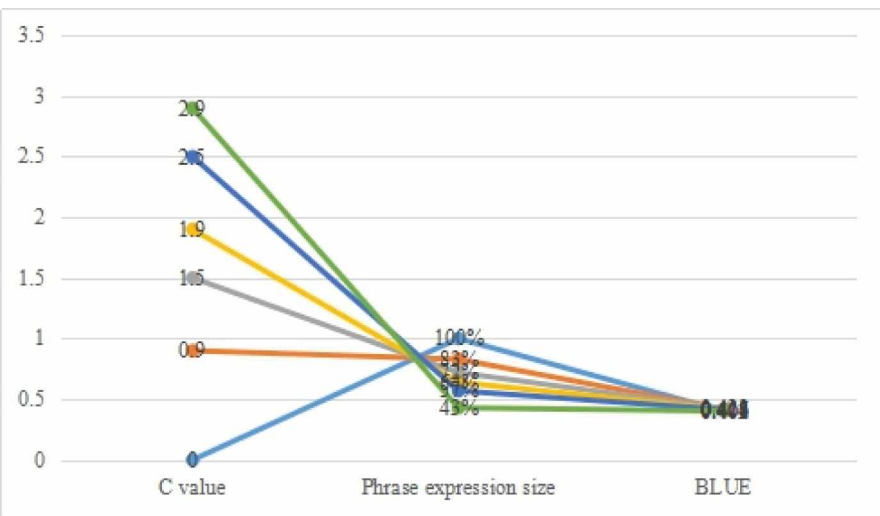


Figure 8. The influence of C value adhesion on Translation

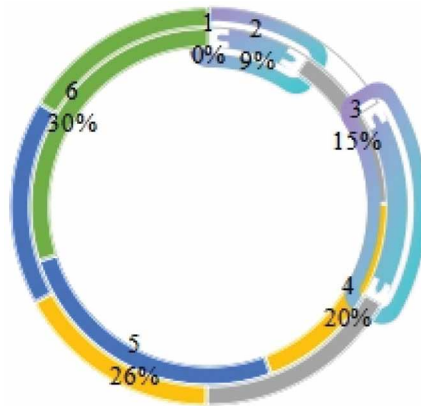


Table 4. Cohesion and Translation

Theme	Threshold	Phrase expression size	BLUE	
			Original phrase list	After filtration
1	0.152	68%	0.313	0.324
2	0.135	73%	0.317	0.354
3	0.113	47%	0.322	0.336
4	0.127	52%	0.324	0.341

Figure 9. Application of parallel algorithm

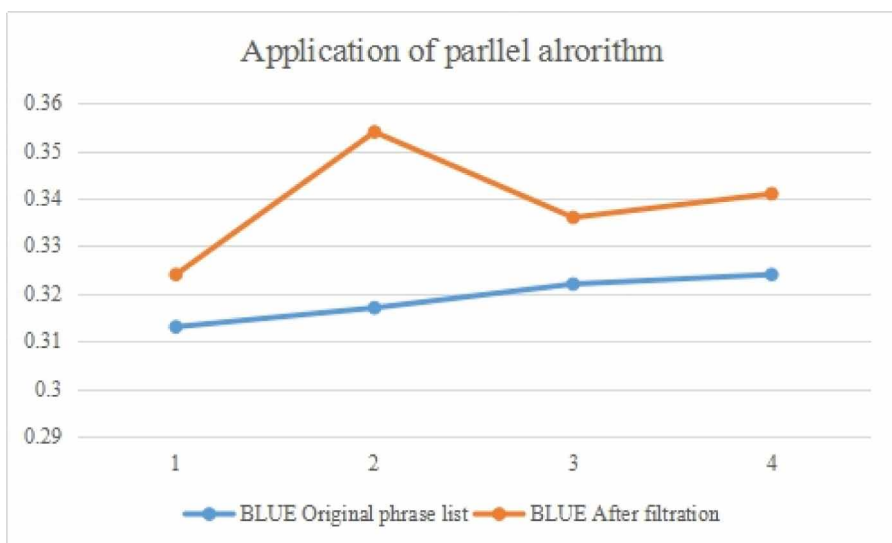
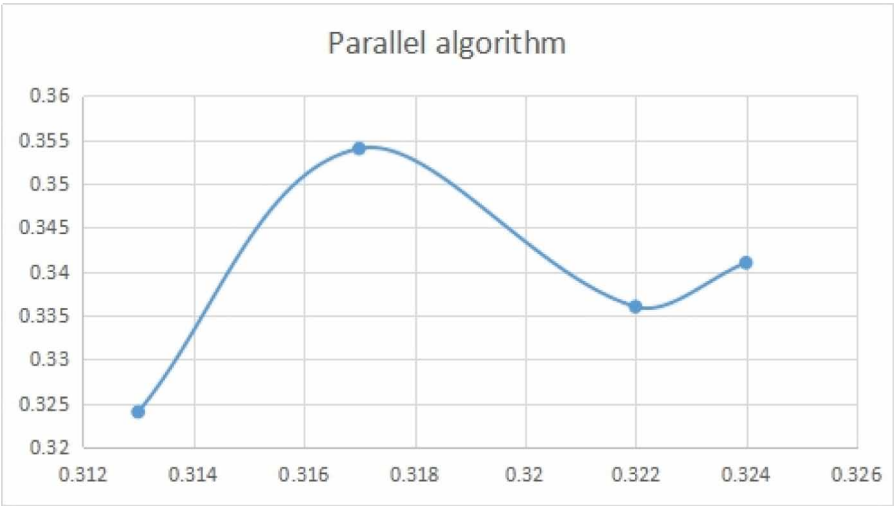


Figure 10. The effect of parallel algorithm on Translation



5. CONCLUSION

Hierarchical phrase translation model uses rules with hierarchical structure as the carrier of translation knowledge, which can capture long-distance ordering. Generally speaking, a translation rule is composed of source language side and target language side, and the two ends can be words, phrases or even syntax trees. The rules of hierarchical phrase translation model are formal syntactic rules, which are extracted from bilingual aligned corpus, and then the eigenvalues of each rule are calculated by maximum likelihood estimation.

The use and selection of translation rules are completely determined by the score obtained during training, without considering the context information and some syntactic information. Firstly, we analyze the translated sentences to get the parsing tree, and then we get whether the rules meet the linguistic information. Finally, we add new features to the model by using syntactic soft constraints, so as to guide the decoding with linguistic information. However, this method gives the same treatment to all rules in the same span, that is, the rules in the same span are not differentiated.

By analyzing the syntax of the training corpus, we extract some features that can represent the context information, train a maximum entropy model, and add this model to the framework of logarithmic linear model, so that the trained maximum entropy model can be used to constrain the translation in the process of decoding. The difference between this method and the previous one is that it does not get the syntactic information from the tested sentences, but directly obtains the relevant information from the training corpus. The maximum entropy model supports multiple features, which provides us with great convenience.

REFERENCES

- Bhadoria, R. S., & Chaudhari, N. S. (2019). Pragmatic Sensory Data Semantics with Service-Oriented Computing. *Journal of Organizational and End User Computing*, 31(2), 22-36.
- Biswas, S., Devi, D., & Chakraborty, M. (2018). A Hybrid Case Based Reasoning Model for Classification in Internet of Things (Iot) Environment. *Journal of Organizational and End User Computing*, 30(4), 104-122.
- Buckley, P., Noonan, S., Geary, C., Mackessy, T., & Nagle, E. (2019). An Empirical Study of Gamification Frameworks. *Journal of Organizational and End User Computing*, 31(1), 22-38.
- Chen. (2019). *Research on Chinese problem generation method of multi encoder neural network based on knowledge map*. Academic Press.
- Deng & Yang. (2016). OS-ELM parallel algorithm based on spark. *Journal of Xi'an University of Posts and Telecommunications*, 21(2), 105-108 + 122.
- Fang, Xu, & Wang. (2017). Japanese English hierarchical phrase translation model based on temporal features. *Computer and Modernization*, 1(6), 1-7.
- Feng, Z. (2015). Statistical machine translation based on phrase and syntax. *Journal of Yanshan University*, 39(6), 546-555.
- Gulen, , Wang, , & Si, . (2016). Post processing of translated verbs in Chinese Mongolian machine translation. Chinese. *Journal of Information Technology*, 30(2), 213-216.
- . Hu, , Wang, , & Song, . (2018). Axial extraction method of pore network based on complete parallel thinning algorithm. *Acta Chem Sinica*, 69(S1), 33-40.
- Kitouni, I., Benmerzoug, D., & Lezzar, F. (2018). Smart Agricultural Enterprise System Based on Integration of Internet of Things and Agent Technology. *Journal of Organizational and End User Computing*, 30(4), 64-82.
- Li, Liang, & Sun. (2017). A machine translation model incorporating the largest bilingual noun phrase. *Computer Application Research*, 34(5), 1316-1320.
- Liu, Zhao, & Cao. (2016). Hierarchical segmentation model based on Markov random field in hierarchical phrase translation. *Journal*, 23(12), 112-115.
- Liu, Y., Qiao, X., & Zhao, S. (2017). Deep fusion of large-scale features in statistical machine translation. *Journal of Zhejiang University*, 1(14), 135-137.
- Lu, Na, & Zhang. (2019). Neural machine translation method based on phrase knowledge. *Chinese Information*, 1(10), 200-201.
- Ma, Ying, & Qiang. (2018). Particle swarm optimization parallel algorithm for multi-objective reservoir scheduling based on spark. *Journal of Xi'an University of Technology*, 34(3), 309-313.
- Namasudra, S., & Roy, P. (2018). Ppbac: Popularity Based Access Control Model for Cloud Computing. *Journal of Organizational and End User Computing*, 30(4), 14-31.
- Qi. (2018). Analysis of mistranslation of interlanguage phrases in machine translation. *Overseas English*, 374(10), 224-225.
- Qiang, L., Han, Y., & Tong, X. (2018). Data generalization and phrase generation in neural machine translation. Chinese. *Journal of Information Technology*, 32(8), 42-52.
- Rong, Z. B. (2015). Implementation of Hadoop based parallel algorithm and analysis of GPS data example. *Southwest University*, 1(42), 153-155.
- Rui, Z. (2017). Design of statistical machine translation system based on phrase similarity. *Zidonghua Yu Yibiao*, 1(8), 66-67, 70.
- Song & Huang. (2017). A Chinese English statistical machine translation method based on syntactic phrases. *Minicomputer System*, 14(10), 235-237.

- Sun, J., Zhu, J., & Yang, F. (2019). Parallel algorithm for Markov clustering of large scale biological networks. *Computer Applications (Nottingham)*, 39(1), 66–71.
- Sun, S., Peng, D., & Huang, D. (2015). Using syntactic phrases to improve the performance of statistical machine translation. Chinese. *Journal of Information Technology*, 29(2), 95–102.
- Tong & Zhu. (2016). Hierarchical phrase machine translation decoding method based on tree to string model reinforcement. *Acta Computer Sinica*, 39(4), 808-821
- Wan, Yu, & Wu. (2015). Tibetan phrase syntax for machine translation. *Computer Engineering and Applications*, 1(13), 211-215+250.
- Wang, Yang, & Chen. (2015). Parallel optimization algorithm for Chinese language model based on cyclic neural network. *Journal of Applied Sciences*, 33(3), 253-261.
- Wu, Xu, & Zhang. (2015). Research on Chinese Japanese joint word segmentation for phrase statistical machine translation. *Computer Engineering and Application*, 3(18), 142-145.
- Yang, Wang, & Yang. (2019). Research and implementation of phrase based Chinese Uyghur machine translation decoding. *Computer Engineering and Design*, 40(4), 1183-1189.
- Yang, Zhengtao, & Ke. (2017). Research on hierarchical phrase translation method integrating linguistic features in lexicalized ordering model. *Computer and Digital Engineering*, 45(12), 2389-2392.
- Zhai, D., Hou, J., & Yue, L. (2019). Parallel algorithm for text sentiment analysis based on deep learning. *Journal of Southwest Jiaotong University*, 3(15), 130–131.
- Zhang. (2015). Research on machine translation method of whole patent document based on phrase recognition. *Informatization Construction*, 204(8), 87 + 89.
- Zhang. (2020). Design of automatic English translation system based on machine intelligent translation and phrase translation. *Automation Technology and Application*, 39(300), 58-61.
- Zhang & Zhang. (2015). Application of phrase extraction algorithm in phrase statistical machine translation. *Heilongjiang Science and Technology Information*, 27(1), 1279-1281.
- Zhu, Yang, & Mi. (2017). Corpus selection technology for Uyghur Chinese machine translation based on bilingual sentence pair coverage. *Journal of the University of Science and Technology of China*, 47(4), 283-289.