

Wuhan University Published Two Duplicated Articles

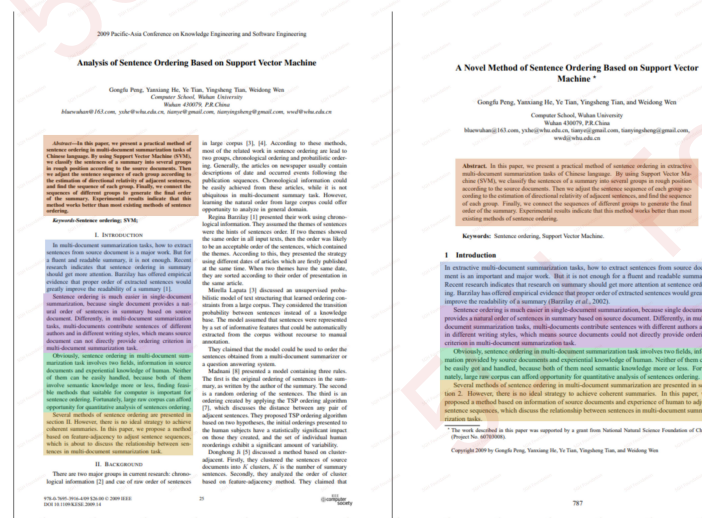
Recently, two articles [1-2] by authors from Wuhan University are found to be more similar than expected.

Article [1], titled "Analysis of Sentence Ordering Based on Support Vector Machine", were published on proceeding of "2009 Pacific-Asia Conference on Knowledge Engineering and Software Engineering".

Article [2], titled "A Novel Method of Sentence Ordering Based on Support Vector Machine", were published on "Proceedings of the 23rd Pacific Asia Conference on Language, Information and Computation"

Both of article [1] and [2] were authored by a same group from Wuhan University: Gongfu Peng, Yanxiang He, Ye Tian, Yingsheng Tian, Weidong Wen (文卫东).

Although the article [2] provides some but not much additional method details and analysis on the results, most of the two articles [1-2] is overlapped:



2009 Pacific Asia Conference on Knowledge Engineering and Software Engineering

Analysis of Sentence Ordering Based on Support Vector Machine

Gangfu Peng, Yanning He, Xu Tian, Yongfang Tian, Weidong Yuan
Computer School, Wuhan University
Wuhan 430079, P.R.China
hwpeng@whu.edu.cn, yhe@whu.edu.cn, tianxu@whu.edu.cn, tyf@whu.edu.cn

Abstract—In this paper, we present a practical method of sentence ordering in multi-document summarization task in Chinese language. By using Support Vector Machine (SVM), we classify the sentences of each source document into two groups according to the nature of sentences. Then, we select the sentences of each group according to the number of sentences in each group. Finally, we return the sentences of different groups to generate the final order of the sentences. Experimental results indicate that this method works better than other existing methods of sentence ordering.

Keywords—sentence ordering; SVM

I. INTRODUCTION

In multi-document summarization task, how to extract sentences from source document is a major work. But for a fixed text and source document, it is not enough. Recent research indicates that sentence ordering in summary should get more attention. Recently, few different empirical evidence that proper order of selected sentences would greatly improve the readability of a summary [1].

Sentence ordering is much easier in single-document summarization, because single document provides a natural order of sentences in summary based on source document. Differently, in multi-document summarization task, multi-document combines sentences of different sources and no difference among them, which means source documents are not in any particular order. Therefore, sentence ordering in multi-document summarization task

becomes a non-trivial problem. In this paper, we present a practical method of sentence ordering in multi-document summarization task. In this method, we use Support Vector Machine (SVM) to classify the sentences of each source document into two groups according to the nature of sentences. Then, we select the sentences of each group according to the number of sentences in each group. Finally, we return the sentences of different groups to generate the final order of the sentences. Experimental results indicate that this method works better than other existing methods of sentence ordering.

Keywords—sentence ordering; SVM

I. INTRODUCTION

In multi-document summarization task, how to extract sentences from source document is a major work. But for a fixed text and source document, it is not enough. Recent research indicates that sentence ordering in summary should get more attention. Recently, few different empirical evidence that proper order of selected sentences would greatly improve the readability of a summary [1].

Sentence ordering is much easier in single-document summarization, because single document provides a natural order of sentences in summary based on source document. Differently, in multi-document summarization task, multi-document combines sentences of different sources and no difference among them, which means source documents are not in any particular order. Therefore, sentence ordering in multi-document summarization task

becomes a non-trivial problem. In this paper, we present a practical method of sentence ordering in multi-document summarization task. In this method, we use Support Vector Machine (SVM) to classify the sentences of each source document into two groups according to the nature of sentences. Then, we select the sentences of each group according to the number of sentences in each group. Finally, we return the sentences of different groups to generate the final order of the sentences. Experimental results indicate that this method works better than other existing methods of sentence ordering.

Keywords—sentence ordering; SVM

TABLE I
ACCURACY OF CLASSIFICATION WITH VARIATION IN THE NUMBER OF SENTENCES

Number of sentences	Accuracy
10	0.85
20	0.88
30	0.90
40	0.92
50	0.94
60	0.96
70	0.98
80	0.99
90	1.00

Our model had solved the problem of noise elimination required by the feature adjacency based ordering.

III. IMPROVED ALGORITHM

Source documents in multi-document summarization task can not provide order information directly. Being relevant parts of text, source documents definitely contain the clue of order, because each of them describes some aspects of the same topic. Thus, we can use the clue of order to generate the order of sentences in source document, which can be the reference standard of sentence ordering.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

TABLE I
ACCURACY OF CLASSIFICATION WITH VARIATION IN THE NUMBER OF SENTENCES

Number of sentences	Accuracy
10	0.85
20	0.88
30	0.90
40	0.92
50	0.94
60	0.96
70	0.98
80	0.99
90	1.00

Our model had solved the problem of noise elimination required by the feature adjacency based ordering.

III. IMPROVED ALGORITHM

Source documents in multi-document summarization task can not provide order information directly. Being relevant parts of text, source documents definitely contain the clue of order, because each of them describes some aspects of the same topic. Thus, we can use the clue of order to generate the order of sentences in source document, which can be the reference standard of sentence ordering.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

2 Related Work

There are two major groups in current research: chronological information [Shankar et al., 2004] and use of time order of sentences in large corpus [Lapin, 2003; Barzilay and Lee, 2004]. Also, the methods for sentence ordering are divided into two groups: chronological ordering and probabilistic ordering. Generally, the articles on newspaper usually contain descriptions of date and events following the publication sequences. Chronological information could be easily achieved from these articles, while it is not ubiquitous in multi-document summary task. However, learning the natural order from large corpus could offer opportunity to analyze sequences in general domain.

Barzilay (Barzilay et al., 2002) presented their work using chronological information. They assumed the themes of sentences were the basis of sentences order. According to this, they presented the strategy by using the dates of different articles, which were firstly published as the order of the sentences. When two themes have the same date, they are sorted according to their order of presentation in the same article.

Michail Lapin (Lapin, 2003) discussed an unsupervised probabilistic model of text structuring that learned ordering constraints from a large corpus. They considered the transition probability between sentences instead of a knowledge base. The model assumed that sentences were represented by a set of informative features that could be automatically extracted from the corpus without recourse to manual annotations.

They claimed that the model could be used to order the sentences obtained from a multi-document summarization or a question answering system.

Michail Lapin (Lapin, 2003) discussed an unsupervised probabilistic model of text structuring that learned ordering constraints from a large corpus. They considered the transition probability between sentences instead of a knowledge base. The model assumed that sentences were represented by a set of informative features that could be automatically extracted from the corpus without recourse to manual annotations.

They claimed that the model could be used to order the sentences obtained from a multi-document summarization or a question answering system.

Michail Lapin (Lapin, 2003) discussed an unsupervised probabilistic model of text structuring that learned ordering constraints from a large corpus. They considered the transition probability between sentences instead of a knowledge base. The model assumed that sentences were represented by a set of informative features that could be automatically extracted from the corpus without recourse to manual annotations.

They claimed that the model could be used to order the sentences obtained from a multi-document summarization or a question answering system.

2 Related Work

There are two major groups in current research: chronological information [Shankar et al., 2004] and use of time order of sentences in large corpus [Lapin, 2003; Barzilay and Lee, 2004]. Also, the methods for sentence ordering are divided into two groups: chronological ordering and probabilistic ordering. Generally, the articles on newspaper usually contain descriptions of date and events following the publication sequences. Chronological information could be easily achieved from these articles, while it is not ubiquitous in multi-document summary task. However, learning the natural order from large corpus could offer opportunity to analyze sequences in general domain.

Barzilay (Barzilay et al., 2002) presented their work using chronological information. They assumed the themes of sentences were the basis of sentences order. According to this, they presented the strategy by using the dates of different articles, which were firstly published as the order of the sentences. When two themes have the same date, they are sorted according to their order of presentation in the same article.

Michail Lapin (Lapin, 2003) discussed an unsupervised probabilistic model of text structuring that learned ordering constraints from a large corpus. They considered the transition probability between sentences instead of a knowledge base. The model assumed that sentences were represented by a set of informative features that could be automatically extracted from the corpus without recourse to manual annotations.

They claimed that the model could be used to order the sentences obtained from a multi-document summarization or a question answering system.

Michail Lapin (Lapin, 2003) discussed an unsupervised probabilistic model of text structuring that learned ordering constraints from a large corpus. They considered the transition probability between sentences instead of a knowledge base. The model assumed that sentences were represented by a set of informative features that could be automatically extracted from the corpus without recourse to manual annotations.

They claimed that the model could be used to order the sentences obtained from a multi-document summarization or a question answering system.

Michail Lapin (Lapin, 2003) discussed an unsupervised probabilistic model of text structuring that learned ordering constraints from a large corpus. They considered the transition probability between sentences instead of a knowledge base. The model assumed that sentences were represented by a set of informative features that could be automatically extracted from the corpus without recourse to manual annotations.

They claimed that the model could be used to order the sentences obtained from a multi-document summarization or a question answering system.

TABLE I
ACCURACY OF CLASSIFICATION WITH VARIATION IN THE NUMBER OF SENTENCES

Number of sentences	Accuracy
10	0.85
20	0.88
30	0.90
40	0.92
50	0.94
60	0.96
70	0.98
80	0.99
90	1.00

Our model had solved the problem of noise elimination required by the feature adjacency based ordering.

III. IMPROVED ALGORITHM

Source documents in multi-document summarization task can not provide order information directly. Being relevant parts of text, source documents definitely contain the clue of order, because each of them describes some aspects of the same topic. Thus, we can use the clue of order to generate the order of sentences in source document, which can be the reference standard of sentence ordering.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

TABLE I
ACCURACY OF CLASSIFICATION WITH VARIATION IN THE NUMBER OF SENTENCES

Number of sentences	Accuracy
10	0.85
20	0.88
30	0.90
40	0.92
50	0.94
60	0.96
70	0.98
80	0.99
90	1.00

Our model had solved the problem of noise elimination required by the feature adjacency based ordering.

III. IMPROVED ALGORITHM

Source documents in multi-document summarization task can not provide order information directly. Being relevant parts of text, source documents definitely contain the clue of order, because each of them describes some aspects of the same topic. Thus, we can use the clue of order to generate the order of sentences in source document, which can be the reference standard of sentence ordering.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

In order to generate the order of sentences in source document, we need to classify the sentences of each source document into two groups. Firstly, we train the model of classification with the positive information of representative sentences in source documents. Secondly, we predict the sentence position in summary. Support Vector Machine (SVM) is a kind of supervised learning method for classification, and we use LIBSVM as classification tool in our model [12].

We gather the first sentence of each paragraph, and put them into training set. For a sentence of summary which is already in training set, we just simply remove it. The label S_{ij} is calculated as $S_{ij} = 0$, where i is the sequence number of the selected sentence in the source document, and j is the number of all sentences in source document. For document D contains 10 sentences, in 3 paragraphs, and each sentence 10 sentences. In this case, $N = 30$, and 3 sentences are selected, which are $s_1 = 1, s_2 = 11, s_3 = 21$, respectively.

709

