

Enhancing Document-Level Relation Extraction with Relation-Specific Entity Representation and Evidence Sentence Augmentation

Qizhu Dai^a, Jiang Zhong^{a,*}, Wei Zhu^a, Chen Wang^a, Hong Yin^a, Qin Lei^a, Xue Li^b and Rongzhen Li^a

^aCollege of Computer Science, Chongqing University, Chongqing, 400044, PR China

^bSchool of Information Technology and Electronic Engineering, The University of Queensland, Brisbane, QLD 4072, Australia

ORCID ID: Qizhu Dai <https://orcid.org/0000-0003-1072-7847>, Jiang Zhong <https://orcid.org/0000-0002-5169-4634>

Abstract. Document-level relation extraction (DocRE) is an important task in natural language processing, with applications in knowledge graph construction, question answering, and biomedical text analysis. However, existing approaches to DocRE have limitations in predicting relations between entities using fixed entity representations, which can lead to inaccurate results. In this paper, we propose a novel DocRE model that addresses these limitations by using a relation-specific entity representation method and evidence sentence augmentation. Our model uses evidence sentence augmentation to identify top-k evidence sentences for each relation and a relation-specific entity representation method that aggregates the importance of entity mentions using an attention mechanism. These two components work together to capture the context of each entity mention in relation to the specific relation being predicted and select evidence sentences that support accurate relation identification. Finally, we re-predicts entity relations based on the evidence sentences, called relationship reordering module. This module re-predicts entity relationships based on the predicted set of evidence sentences to form k sets of relationship predictions, and then averages these k+1 sets of results to obtain the final relationship predictions. Experimental results on the DocRED dataset demonstrate that our proposed model achieves an F1 score of 62.84% and an Ign F1 score of 60.79%, outperforming state-of-the-art methods.

1 INTRODUCTION

Relation extraction (RE) is a very important task in natural information extraction, aiming at identifying the semantic relationships between entities in a given text. It has rich applications in knowledge graph construction, question and answer, and biomedical text analysis. applications[18, 10, 2, 16]. Previous studies have mainly focused on predicting the relationship between two mentions in a sentence. However, in practice an entity may have multiple entity referents throughout a document, and relationships can only be inferred given multiple sentences as contexts. The relationship can only be inferred given multiple sentences as contexts. Hence, due to its potential practical applications, research on document-level relation extraction (DocRE) has garnered significant attention in recent years[21, 22, 16].

Unlike document-level relationship extraction studies that take the entire document as input, humans may only need a few sentences to infer the relation of entity pairs. As can be seen in Figure 1, only sentence S1 is needed to determine that *SamuelHerbertCohen* is a citizen of Australia, and the same birthplace and school can be identified by sentence S2 alone. According to Huang et.al.[5], over 95% of the instances in the DocRED[16] dataset required no more than three sentences as supporting evidence, but the average document in this dataset had eight sentences each. When using all sentences in a whole document for relationship classification, there are inevitably long-distance dependencies between entities.

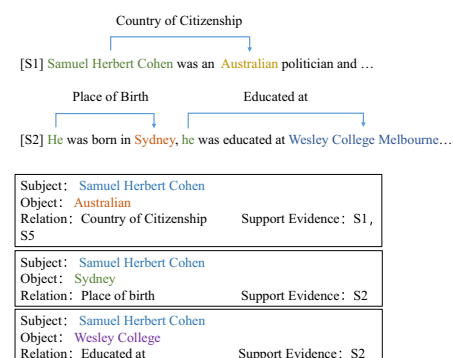


Figure 1. A test sample in the DocRED dataset. The mention of the entity *SamuelHerbertCohen* in sentence S1 is more important for the classifier to identify the relation between this entity and *Australian*. However, the second referent *He* should be taken into account more in order to identify the relation *Placeofbirth*. This suggests that different entity mentions should play different roles in the identification of different relationships involving the same entity.

Therefore, by identifying the evidence sentences that contribute significantly to the entity for relation extraction, it can be of great help in determining the relation of this entity pair, and the identified evidence sentences can be used as a model to predict the interpretation of the entity pair as a certain relationship, increasing the explanatory power of the model. Yao et al.[16] introduced the evidence

* Corresponding Author. Email: zhongjiang@cqu.edu.cn

extraction task, and Huang et al.[3] used the evidence extraction task as a secondary task to enhance the relationship extraction effect of the model. The EIDER model [13] first trains a relation extraction (RE) model jointly with a lightweight evidence extraction model that is efficient in terms of both memory and runtime. However, these works only predict whether each sentence is evidence or not separately based on the given entity pairs, and do not use the extracted evidence sentences in the subsequent process.

Furthermore, a key step in document-level relationship extraction is to obtain a representation of entities from their individual denotations. In previous studies, relationship extraction models have simply applied average pooling [17, 14], and maximum pooling [8] (or Log-SumExp pooling [22, 21]) to compute a fixed representation for a given entity, which is then fed into a classifier for relationship classification. However, different entity referents of entities in a document may have different semantics, and simply generating a fixed entity representation may confuse the semantics of different entity referents and reduce the accuracy of the relationship classification of an entity when it has multiple relationship instances, which has not been considered in previous studies.

To address the above issues, this paper proposes a document-level relationship extraction model with adaptive entity denotation representation and evidence sentence augmentation. Firstly, we propose a relation-specific entity representation method for relation prediction. The method requires training a representation under each relationship aspect, then using an attention mechanism to calculate the importance of entity mentions, and finally aggregating the entity mention representations under different relations based on the attention scores to obtain a relation-specific entity representation. This representation is then fed into a bilinear network to determine whether the entity pair has a certain relation. In the evidence extraction module, this paper uses the relation-specific entity representation to calculate whether sentence j is a supporting evidence sentence for the entity pair (e_h, e_t) prediction relation r . Specifically, we sort the relations predicted in the initial relation prediction module according to their probability magnitude, and then select the top- k relations among them for the prediction of evidence sentences, eventually forming a collection of evidence sentences for each relation. Finally, the relationship reordering module re-predicts the entity relationships based on the predicted set of evidence sentences to form k sets of relationship predictions and then averages these $k+1$ sets of results to obtain the final relationship prediction results. The experimental results show that all the proposed methods in this thesis have significantly improved the accuracy of relation extraction.

CONTRIBUTIONS. (1)To overcome the limitations of existing DocRE models that use fixed entity representations, we present a novel approach that utilizes a relation-specific entity representation method and incorporates evidence sentence augmentation. (2) To improve the accuracy of our proposed DocRE model, we incorporated a relation reordering module that leverages evidence sentences obtained from the evidence extraction task. This module addresses issues such as incomplete extraction of evidence sentences and missing information. (3)We evaluate our model to achieve state-of-the-art performance on the DocRED dataset.

2 RELATED WORK

Document-level relation extraction(DocRE) is a challenging task that has received significant attention in recent years [6, 7, 16]. Prior work has explored various approaches for performing document-level relation extraction, including both Graph-based and Transformer-based

methods.

2.1 Graph-based DocRE

Graph-based methods for document-level relation extraction typically construct a graph with nodes representing mentions, entities, sentences, or documents, and use graph-based reasoning to infer relations. Zeng et al.[20]perform multi-hop reasoning on both mention-level and entity-level graphs. Xu et al. [15] extract reasoning paths for each relation and train the model to reconstruct these paths. Zeng et al. [19] separately handle intra- and inter-sentential entity pairs and perform multi-hop reasoning on a mention-level graph for inter-sentential entity pairs. However, constructing a graph may result in important information being omitted from the text, and complex operations on the graphs may hinder the model's ability to capture the text structure.

2.2 Transformer-based DocRE

In addition to the graph-based approach, another direction of research is to model inter-sentence associations through the implicit capture of long-range inter-token dependency using a transformer architecture[5]. Zhang et al.[21] treat document-level relationship extraction as a semantic segmentation task on the entity matrix and apply U-Net to capture the correlations between relationships. Zhou et al.[22] use an attention mechanism in the transformer to extract useful context, and an adaptive threshold is applied to each entity pair. Huang et al.[4]extract evidence at the document level to guide the discovery of relations. However, this approach has a significant runtime and memory overhead, as it is highly dependent on the evidence annotations. In contrast, Huang et al. [5] predict only a few rule-selected sentences, which may miss important information and does not consistently improve performance. In contrast to these approaches, the EIDER model [13] includes a lightweight evidence extraction model that is substantially faster than Huang et al. [4] and improves relation extraction at the document level even when trained on gold tags. By effectively extracting evidence and integrating it into reasoning, EIDER model can enhance its capabilities.

3 Methodology

3.1 Problem Formulation

The task of document-level relation extraction (DocRE) is as follows: given a document D consisting of N sentences $(\{S_n\}_{n=1}^N)$, L tokens $(\{h_l\}_{l=1}^L)$, E named entities $(\{e_i\}_{i=1}^E)$ and the mentions (m_j^i) of each entity, the task of DocRE is to predict the possible relations between all entity pairs (e_h, e_t) from a predefined set of relations $(R \cup \{NA\})$, with e_h and e_t representing the head entity and tail entity, respectively. The extracted set of evidence sentences are combined together in their order in the document to form a new input D_{evi} , and the entity pairs are added to the front of the document to form (e_h, e_t, D_{evi}) , which is fed into the relational reordering model to calculate the probability value of the relationship between the entity pairs (e_h, e_t) .

3.2 Method

In this paper, BERT is used as the base encoder, and as can be seen in Figure 2, the model is divided into three main parts, namely the relationship prediction module, the evidence extraction module, and the relationship reordering module.

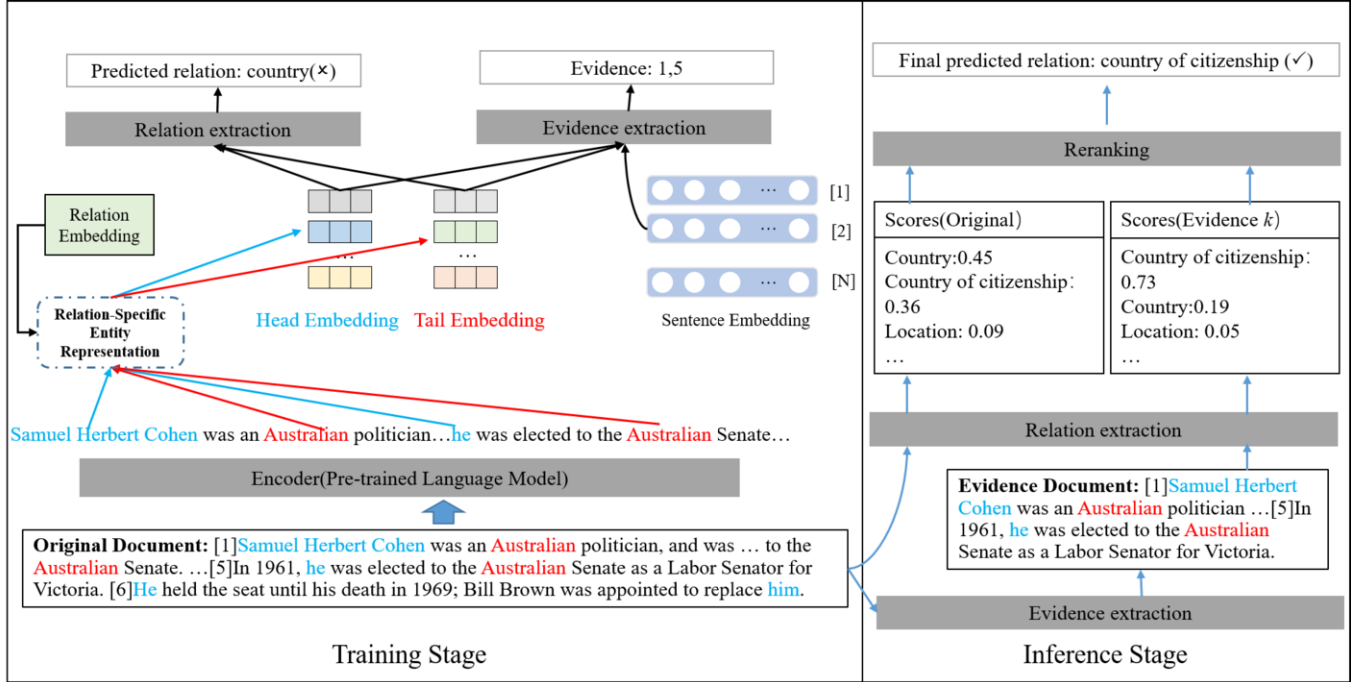


Figure 2. An illustration of the framework. The left part illustrates the training stage and the right shows the inference stage. Relation-specific entity representation module is a component in the proposed DocRE model. It can aggregates the importance of entity mentions using an attention mechanism to create a relation-specific representation of each entity. The evidence extraction module is responsible for identifying evidence sentences that support accurate relation identification. And The relation reordering module is responsible for re-predicting entity relations based on the evidence sentences identified by the evidence extraction module.

Base Encoder. We encode the semantics of each token in the document using a pre-trained language model [1]. Given a document $D = \{x_n\}_{n=1}^L$, the special tokens $\langle STA \rangle$ and $\langle END \rangle$ are inserted before the start position and after the end position of the entity mention designation to mark the location of the entity mention m_j^i .

$$\begin{aligned} H &= [h_0, h_1, h_2, \dots, h_i, \dots, h_{L+3}] \\ &= \text{Encoder}([\langle STA \rangle, x_1, x_2, \langle END \rangle, \dots, X_L])(2) \end{aligned}$$

Where h_i is the vector representation of $token_i$. We utilize the embedding of the special token $\langle STA \rangle$ to the mention of this entity. Then, previous works obtained the embedding of entity e_i by applying LogSumExp pooling [6, 22] to the embedding of all its mentions.

$$e_i = \frac{1}{Q_i} \sum_{j=1}^{Q_i} m_j^i \quad (3)$$

Where m_j^i denotes the representation of the j th mention of the entity e_i .

However, when considering different relations, using a fixed representation may overlook the different contributions of different mentions of entities to a particular relation r , so this paper proposes a relation-specific entity representation, as shown in Fig 3.

Specially, all relations learn the corresponding representation R_r , which is obtained by random initialization and subsequently updated during the training of the model. The semantic relatedness between relations and mention is then calculated using the following equation.

$$S_{ij}^r = f(R_r, m_j^i) \quad (4)$$

Where R_r is the vector representation of relation r and f denotes the function that calculates the similarity between the vector representation of r and the j th mention of entity e_i . The dot product approach

is finally chosen in this paper. Next, we obtain the final relational attention score α_{ij}^r using the normalized exponential function Softmax for the semantic relational relevance score S_{ij}^r for all the mentions of entity e_i to the relation r . Ultimately, this paper uses the addition with weights for all alleged vector representations of entities to obtain the entity representation of the relational characteristics e_i^r .

$$\alpha_{ij}^r = \frac{S_{ij}^r}{\sum_{k=1}^n S_{ik}^r} \quad (5)$$

$$e_i^r = \sum_{j=1}^n \alpha_{ij}^r m_j^i \quad (6)$$

Where m_j^i denotes the representation of the j th mention of the entity e_i . n is the number of mentions corresponding to the entity e_i .

Relation Classifier. First, the representation e^r of an entity over a relation r is obtained according to the proposed relation-specific entity representation, after which the entity pair (e_h^r, e_t^r) is directly fed into a bilinear layer to calculate the probability that the relation between the entity pairs is r .

$$P(r | e_h^r, e_t^r) = \sigma(e_h^r W_r e_t^r + b_r) \quad (7)$$

where W_r and b_r are trainable parameters. Here we utilize cross entropy as a loss function for the relation classifier.

$$\mathcal{L}_{re} = - \sum_{(h,t) \in D} \sum_{r \in P} \tilde{y}_r(e_h^r, e_t^r) \log(P(r | e_h^r, e_t^r)) \quad (8)$$

where $\tilde{y}_r \in \{0, 1\}$, $\tilde{y}_r = 1$ denotes the relation of the entity pair (e_h, e_t) is r .

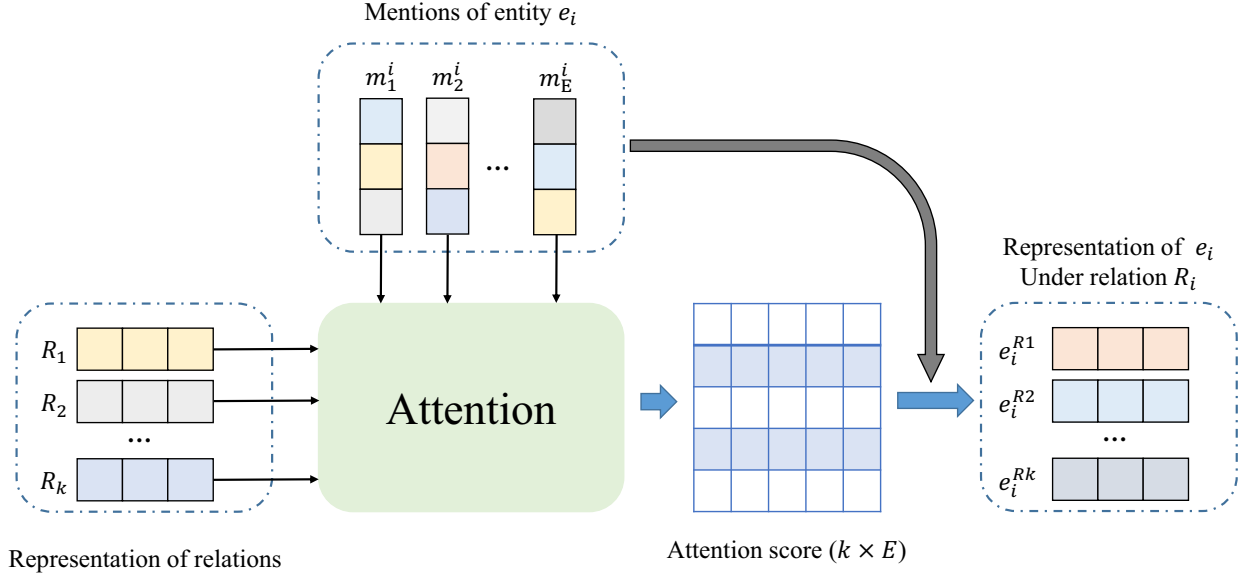


Figure 3. Relation-Specific Entity Representation.

Evidence Classifier. Evidence sentences contain important information for predicting relations between head and tail entities. In the Relation Classifier task, it is necessary to extract entities from text and identify relationships between them, the model may ignore some important contextual information when performing the Relation Classifier task alone, which may lead to incorrect predictions. Therefore, by predicting evidence sentences, the model can better understand the textual context and can be more accurate in inferring relations between entities.

Therefore, we also predict whether each sentence S_n is an evidence sentence of entity pair (e_h^r, e_t^r) . First, a set of prediction results for the entity pair (e_h^r, e_t^r) is obtained by the first stage of relation extraction, and the top-k relations in the results are selected for the evidence sentence extraction. Then, the representation of the sentence $S_n = \log \sum_{token_i \in S_n} \exp(token_i)$ by applying LogSumExp pooling [6, 22] on all tokens of the sentence S_n . If sentence S_n is the evidence sentence of entity pair (e_h^r, e_t^r) , the token in S_n may be relevant to the relation classifier and contribute more to the representation of entities.

$$P(S_n | e_h^r, e_t^r) = \sigma(S_n W_v [e_h^r \oplus e_t^r] + b_v) \quad (9)$$

where $\oplus$ denotes splicing two vectors. W_v and b_v are learnable parameters. Since an entity pair may have multiple evidence sentences, we use the cross-entropy loss as the objective function to optimize the model.

$$\mathcal{L}_{Evi} = - \sum_{h \neq t} \sum_{S_n \in \mathcal{D}} y_n P(S_n | e_h^r, e_t^r) \quad (10)$$

$$+ (1 - y_n) \log(1 - P(S_n | e_h^r, e_t^r)) \quad (11)$$

where $y_n = 1$ indicates that sentence S_n is the evidence sentence for entity pair (e_h^r, e_t^r) for relation r .

Finally, we optimize our model using the joint optimization of the relation extraction loss $\mathcal{L}_{re}$ and evidence extraction loss $\mathcal{L}_{Evi}$:

$$\mathcal{L} = \mathcal{L}_{re} + \mathcal{L}_{Evi} \quad (12)$$

4 Relation Reordering Module for Inference

Although the extracted evidence sentences already contain all the information relevant to the relation, no system can extract evidence perfectly without missing any sentences, and the extracted evidence sentences may be false. Thus relying on extracted evidence alone may miss important information in the document and lead to sub-optimal performance. Therefore, we combine the predictions from the original document and the extracted evidence to reorder and predict. Reordering is a significant step in improving the accuracy of relation identification. The relational reordering module is used to re-predict entity relations based on the evidence sentences identified by the evidence extraction module. The purpose of the reordering step is to refine the initial relation predictions made by the model based on the evidence sentences. By reordering the predicted relations based on the evidence sentences, the model can improve the accuracy of its predictions and ensure that the final output reflects the most relevant and accurate relationships between entities.

Especially, as shown in Figure 2, we first obtain a set of evidence sentences generated using the entity top-k relations. Then, they are fed into the relationship reordering model to calculate the relationship probability values between entity pairs (e_h, e_t) , and the relationship reordering module has the same network structure as the initial relationship prediction module. The k sets of (e_h, e_t, D_{evi}) are fed into the model for relation prediction to obtain k sets of relationship probability values, and finally, we average these k sets of relation-

ship probability values with the initial relationship probability values to obtain the final relationship probability values.

5 Experiments

5.1 Experiment Setup

Datasets. The effectiveness of EIDER is evaluated on DocRED dataset [16]. The DocRED dataset was manually annotated with 5,053 documents from the English Wikipedia, of which 1,000 were used as the validation set and 1,000 as the test set, for a total of 132,375 entities and 56,354 relationship instances, making it the largest manually annotated document set for relationship extraction. At least 40.7% of the relationship instances in the DocRED dataset need to be identified by integrating information from multiple sentences in a document, so the model needs to read multiple sentences in a document to identify entities. Therefore, the model needs to read multiple sentences in the document to identify entities and infer the relationship between entities by combining the information in the document.

Baselines. LSR [9]: the model enhances inter-sentence relational reasoning by automatically inscribing potential document-level graphs using GCNs, and proposes a graph adjustment strategy that enables the model to progressively aggregate relevant information for multi-hop reasoning;

GAIN [20]: this model is used to capture the interactions between different entity referents and entity-level graphs for pooling information about all referents of the same entity by constructing entity referent-level graphs and proposes a path inference mechanism to infer the relationships between entities;

BERT-base [11]: the use of BERT to encode the text, the representation of the average entity referent as the representation of the entity, and the subsequent classification of the relationship;

E2GRE [4]: a model that for the first time includes an evidence extraction task as a secondary task to enhance the effectiveness of the model’s relationship extraction;

ATLOP [22]: the model proposes adaptive thresholding to solve the multi-label problem and uses local context pooling to solve the multi-entity problem;

Eider [13]: this model also incorporates an evidence extraction task and incorporates evidence sentences into the reasoning process.

Evaluation Metrics. Consistent with previous studies[16], we adopt **F1** and **Ign F1** as the main evaluation metrics for relation extraction. Ign F1 measures the F1 score that excludes the relationship shared between the training set and the dev/test set.

Implementation Details. The model in this paper is implemented based on PyTorch and Huggingface’s Transformers [12], and the experiments use the cased-BERT base [1] as the base encoder and the output of BERT are mapped to 200 dimensions using MLP. In this paper, we optimize the model using the AdamW optimizer with an initial learning rate set to 5e-5, a batch size set to 4, and weight decay parameter set to 1e-4. To prevent the model from starting to overfit as the number of training rounds increases, we use warmup and early stopping techniques to periodically check the performance of the model on the validation set during training, and stop the model if the performance on the validation set To be fair, all models are trained on a single RTX A6000 GPU. Table 1 displays the model hyperparameters.

Table 1. Details of hyperparameters used for DocRE task.

Hyperparameters	Value
dimension	200
learning rate	5e-5
warmup_rate	0.06
Optimiser	AdamW
Batch Size	4
Epoch	50

5.2 Main Results

We compare our models with both Graph-based methods and transformer-based methods. The overall performance comparison experimental results of the framework are shown in Table 2, and the bolded font in the table indicates the optimal results of the model. It can be seen that the proposed framework in this paper basically achieves the best results compared to the baseline model. E2GRE improves the results by 5.52 only by adding the auxiliary task of evidence extraction, which proves that the evidence extraction task helps significantly in the document-level relationship extraction task. Overall, the experimental results show that the proposed relationship extraction framework based on significant entity referents and sentences is effective.

The comparison to baseline models shows how the proposed model is better than others. But the improvements over the Eider model are small. Eider model is the current state-of-the-art in document-level relationship extraction. The proposed relation-specific entity representation method aggregates the importance of entity mentions using an attention mechanism. And it can help to demonstrate their awareness of the limitations of current work and their commitment to further advancing the field of DocRE.

Table 2. Relation extraction results on DocRED.

Model	Dev		Test	
	<i>lgnF1</i>	<i>F1</i>	<i>lgnF1</i>	<i>F1</i>
LSR	52.43	59.00	56.97	59.05
GAIN	59.14	61.22	59.00	61.24
BERT-base*	53.03	54.95	52.92	54.89
E2GRE	54.91	55.38	55.22	58.72
ATLOP	59.22	61.09	59.31	61.30
Eider	60.51	62.48	60.42	62.47
Ours	60.87	62.91	60.79	62.84

* indicates that the results are reproduced in this paper.

5.3 Ablation Study

To further validate the performance of the modules in the proposed relational abstraction model, multiple sets of ablation experiments are conducted in this paper. In order to verify the effectiveness of the relationship-specific entity representation module, it is removed from the model and the relationship-specific entity representation module is represented by IM using averaging over entity designations. To verify the effectiveness of the proposed relation reordering in this paper, we remove this relation reordering part and tune the initial relation ordering model to the optimal one by increasing the number of training rounds, and RR denotes the evidence extraction and relation reordering module. We also explore the effect of denotational

disambiguation by removing the HOI denotational disambiguation step and using only the original documents for model training. The experimental results are shown in Table 3.

Table 3. Main ablation experiment results.

Model	Dev		Test	
	<i>lgnF1</i>	<i>F1</i>	<i>lgnF1</i>	<i>F1</i>
Ours	60.87	62.91	60.79	62.84
-IM	59.56	62.91	60.79	62.84
-RR	58.36	60.39	58.28	60.25
-HOI	60.17	62.15	60.06	62.11

From Table 3, we can see that the performance of the model decreases by 1.3 after removing the relationship-specific entity representation module from the model, which illustrates the effectiveness of the relationship-specific entity representation module proposed in this paper. Removing the relationship reordering module causes the performance degradation of models 2.5 to 2.6, which illustrates the effectiveness of the proposed relationship reordering module in this paper. From the results of "-HOI", we can see that the performance of the model decreases from 0.7 to 0.76 when using only the original documents for relationship extraction, which indicates that the denotational disambiguation task can bring improvement to the document-level relationship extraction.

Finally, the relationship-specific entity representation module is added to the BERT-base and LSR models to form two new models, BERT-base-IM and LSR-IM, and the experimental results are shown in Table 4. The experimental results are shown in Table 3.7. From the results in the table, it can be seen that the proposed relationship-specific entity representation method can be used in other models to enhance the performance of the models by 1.1 1.38, which shows the effectiveness of the method and can be added as a plug-in to other models to improve the performance of the models.

Table 4. Ablation experiment results of relation representation module.

Model	Dev		Test	
	<i>lgnF1</i>	<i>F1</i>	<i>lgnF1</i>	<i>F1</i>
BERT-base	53.03	54.95	52.92	54.89
BERT-base-IM	54.19	56.05	54.07	56.01
LSR	52.43	59.00	56.97	59.05
LSR-IM	56.38	60.34	58.35	60.29

6 Conclusions

This paper addresses the fact that most existing models simply aggregate entity denotations to obtain a fixed entity representation, ignoring the contribution of different entity denotations to different relationships. This chapter proposes a relationship-specific entity representation approach, which uses an attention mechanism to enable the model to make full use of the highly relevant entity denotation information of the relationship for relationship classification. To address the problem that most models use all sentences in a document for relationship prediction and that models cannot handle long-range dependencies well, this chapter proposes an entity-related evidence extraction task that enables the model to focus on important sentences.

The framework finally makes full use of the evidence sentences obtained from the evidence extraction task and utilizes a relationship reordering module to enhance the accuracy of the model's final relationship extraction. The experimental results show that the model achieves better results on the DocRE dataset, and the effectiveness of the modules in the model is verified through ablation experiments.

We propose a relation-specific entity representation method, which aggregates the importance of entity mentions using an attention mechanism, and passes it through a bilinear network for relation prediction. Overall, the approach taken by the proposed model addresses the problem of different entity referents having different semantics by generating relation-specific entity representations that capture the context of each entity mention in relation to the specific relation being predicted. And it can help to demonstrate their awareness of the limitations of current work and their commitment to further advancing the field of DocRE.

Acknowledgements

This work is funded in part by the National Natural Science Foundation of China under Grants No.62176029, and in part by the graduate research and innovation foundation of Chongqing, China under Grants No.CYB21063. This work also is supported in part by the Chongqing Technology Innovation and Application Development Special under Grants CSTB2022TIAD-KPX0206. We express our sincere gratitude to the above funding agencies. We also acknowledge the computational support from the Chongqing Artificial Intelligence Innovation Center.

References

- [1] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova, 'Bert: Pre-training of deep bidirectional transformers for language understanding', in *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics*, (2019).
- [2] Bayu Distiawan, Gerhard Weikum, Jianzhong Qi, and Rui Zhang, 'Neural relation extraction for knowledge base enrichment', in *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pp. 229–240, (2019).
- [3] Kevin Huang, Peng Qi, Guangtao Wang, Tengyu Ma, and Jing Huang, 'Entity and evidence guided document-level relation extraction', in *Proceedings of the 6th Workshop on Representation Learning for NLP (RepLanLP-2021)*, pp. 307–315, Online, (August 2021). Association for Computational Linguistics.
- [4] Kevin Huang, Peng Qi, Guangtao Wang, Tengyu Ma, and Jing Huang, 'Entity and evidence guided document-level relation extraction', in *Proceedings of the 6th Workshop on Representation Learning for NLP (RepLanLP-2021)*, pp. 307–315, (2021).
- [5] Quzhe Huang, Shengqi Zhu, Yansong Feng, Yuan Ye, Yuxuan Lai, and Dongyan Zhao, 'Three sentences are all you need: Local path enhanced document relation extraction', in *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing. Stroudsburg, PA: ACL*, pp. 998–1004, (2021).
- [6] Robin Jia, Cliff Wong, and Hoifung Poon, 'Document-level n-ary relation extraction with multiscale representation learning', in *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pp. 3693–3704, (2019).
- [7] Po-Ting Lai and Zhiyong Lu, 'Bert-gt: cross-sentence n-ary relation extraction with bert and graph transformer', *Bioinformatics*, **36**(24), 5678–5685, (2020).
- [8] Jingye Li, Kang Xu, Fei Li, Hao Fei, Yafeng Ren, and Donghong Ji, 'Mrn: A locally and globally mention-based reasoning network for document-level relation extraction', in *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021*, pp. 1359–1370, (2021).

- [9] Guoshun Nan, Zhijiang Guo, Ivan Sekulić, and Wei Lu, ‘Reasoning with latent structure refinement for document-level relation extraction’, in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pp. 1546–1557, (2020).
- [10] Yu Shi, Jiaming Shen, Yuchen Li, Naijing Zhang, Xinwei He, Zhengzhi Lou, Qi Zhu, Matthew Walker, Myunghwan Kim, and Jiawei Han, ‘Discovering hypernymy in text-rich heterogeneous information network by exploiting context granularity’, in *Proceedings of the 28th ACM International Conference on Information and Knowledge Management*, pp. 599–608, (2019).
- [11] Hong Wang, Christfried Focke, Rob Sylvester, Nilesch Mishra, and William Wang, ‘Fine-tune bert for docred with two-step process’, *arXiv preprint arXiv:1909.11898*, (2019).
- [12] Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, et al., ‘Transformers: State-of-the-art natural language processing’, in *Proceedings of the 2020 conference on empirical methods in natural language processing: system demonstrations*, pp. 38–45, (2020).
- [13] Yiqing Xie, Jiaming Shen, Sha Li, Yuning Mao, and Jiawei Han, ‘Eider: Empowering document-level relation extraction with efficient evidence extraction and inference-stage fusion’, in *Findings of the Association for Computational Linguistics: ACL 2022*, pp. 257–268, (2022).
- [14] Benfeng Xu, Quan Wang, Yajuan Lyu, Yong Zhu, and Zhendong Mao, ‘Entity structure within and throughout: Modeling mention dependencies for document-level relation extraction’, in *Proceedings of the AAAI conference on artificial intelligence*, volume 35, pp. 14149–14157, (2021).
- [15] Wang Xu, Kehai Chen, and Tiejun Zhao, ‘Document-level relation extraction with reconstruction’, in *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pp. 14167–14175, (2021).
- [16] Yuan Yao, Deming Ye, Peng Li, Xu Han, Yankai Lin, Zhenghao Liu, Zhiyuan Liu, Lixin Huang, Jie Zhou, and Maosong Sun, ‘Docred: A large-scale document-level relation extraction dataset’, in *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pp. 764–777, (2019).
- [17] Deming Ye, Yankai Lin, Jiaju Du, Zhenghao Liu, Peng Li, Maosong Sun, and Zhiyuan Liu, ‘Coreferential reasoning learning for language representation’, in *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pp. 7170–7186, (2020).
- [18] Mo Yu, Wenpeng Yin, Kazi Saidul Hasan, Cicero dos Santos, Bing Xiang, and Bowen Zhou, ‘Improved neural relation detection for knowledge base question answering’, in *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 571–581, Vancouver, Canada, (July 2017). Association for Computational Linguistics.
- [19] Shuang Zeng, Yuting Wu, and Baobao Chang, ‘Sire: Separate intra- and inter-sentential reasoning for document-level relation extraction’, in *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021*, pp. 524–534, (2021).
- [20] Shuang Zeng, Runxin Xu, Baobao Chang, and Lei Li, ‘Double graph based reasoning for document-level relation extraction’, in *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pp. 1630–1640, (2020).
- [21] Ningyu Zhang, Xiang Chen, Xin Xie, Shumin Deng, Chuanqi Tan, Mosha Chen, Fei Huang, Luo Si, and Huajun Chen, ‘Document-level relation extraction as semantic segmentation’, in *Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence, IJCAI*, (2021).
- [22] Wenxuan Zhou, Kevin Huang, Tengyu Ma, and Jing Huang, ‘Document-level relation extraction with adaptive thresholding and localized context pooling’, in *Proceedings of the AAAI conference on artificial intelligence*, volume 35, pp. 14612–14620, (2021).