

FAST-RE: Feature-Augmented Scientific Text Relation Extraction with Logic-Aware Representation Injection

Chaoran Zhou^{1,3}, Xinyu Li¹, Xin Zhang^{1,4}, and Yang Sisi²

¹ Changchun University of Science and Technology, ChangChun, China
zcr@cust.edu.cn, 2328189353@qq.com, zhangxin@cust.edu.cn

² Institute of Scientific and Technical Information of Jilin, Changchun, China
1046235989@qq.com

³ Changchun University of Science and Technology High-Tech Industry Co., Ltd.,
Changchun, China

⁴ Jilin Provincial Science and Technology Innovation Center for Network Database
Application Software, Changchun, China

Abstract. Scientific text RE is essential for constructing high-quality scientific knowledge graphs. However, existing methods mainly rely on generic semantic representations and often fail to explicitly model domain-specific scientific reasoning patterns, making it difficult to identify implicit relations expressed through causal reasoning, parameter association, experimental verification, and innovation evaluation. To address this issue, we propose FAST-RE, a Feature-Augmented Scientific Text RE model with logic-aware representation injection. FAST-RE first encodes scientific sentences with RoBERTa to obtain contextual semantic representations. A scientific feature augmentation and injection module is then designed to capture multi-dimensional logic cues and inject them into contextual representations. The enhanced representations are further refined by an iterated dilated convolutional network to model global contextual interactions, followed by a subject-first decoding strategy for extracting subject-relation-object triples. By coupling contextual semantic understanding with explicit scientific logic modeling, FAST-RE improves the extraction of implicit relations in complex scientific texts. In addition, we construct a fine-grained scientific RE dataset containing 4,000 sentence-level samples and 6,120 annotated triples across 10 relation types. Experimental results show that FAST-RE outperforms competitive baselines in terms of precision and $F1$ score, while ablation studies confirm the effectiveness of the proposed feature augmentation and logic-aware injection mechanisms.

Keywords: Scientific Knowledge Graph · Subject-Relation-Object Triple Extraction · Logic-Aware Representation Injection · Feature Augmentation

1 Introduction

Scientific literature contains a large amount of structured knowledge about research methods, technical problems, experimental evidence, and scientific findings, among others[13]. Automatically extracting semantic relations from unstructured scientific texts is fundamental to constructing high-quality scientific knowledge graphs, which support applications such as technology trend analysis and scientific decision-making. Therefore, scientific text RE has become an important task in natural language processing and scientific knowledge management.

In recent years, the RE methods have achieved promising performance on general-domain RE tasks through contextual semantic representation learning and improved decoding strategies [2]. However, scientific texts differ from general-domain texts in that relations between scientific entities are often not expressed through surface co-occurrence, but are embedded in scientific logic patterns[15]. These characteristics pose significant challenges for scientific RE. Although pre-trained language models provide powerful contextual representations, they mainly rely on implicit semantic association learning and lack explicit modeling of scientific logic. Subject-first extraction frameworks are suitable for scientific text RE, as scientific discourse often centers on core methods, models, or concepts[21]. However, existing subject-first models still largely depend on generic semantic representations for relation and object decoding, limiting their ability to identify implicit relations in complex scientific argumentation.

To address these issues, we propose FAST-RE, a Feature-Augmented Scientific Text RE model with logic-aware representation injection. FAST-RE integrates contextual semantic representations from RoBERTa with explicit scientific logic cues, including causal reasoning, parameter association, experimental verification, and innovation evaluation. These logic-aware features are dynamically injected into semantic representations and further used for global contextual interaction modeling and subject-relation-object triple extraction. In this way, FAST-RE enhances the recognition of implicit relations in complex scientific texts while maintaining an end-to-end extraction. The main contributions of this paper are summarized as follows:

- We propose FAST-RE, a feature-augmented RE model for scientific texts, which integrates contextual semantic representation learning with explicit scientific logic modeling through logic-aware representation injection.
- We design a scientific feature augmentation and injection module that captures causal reasoning, parameter association, experimental verification, and innovation evaluation cues, and dynamically injects them into contextual representations to improve implicit RE.
- We construct a fine-grained scientific text RE dataset containing 4,000 sentence-level samples and 6,120 annotated triples across 10 scientific relation types, and conduct extensive experiments to verify the effectiveness of FAST-RE.

2 Related Work

Existing resources and systems, such as AI-KG [3], and SciERC [13], have promoted the organization and utilization of scholarly knowledge by representing research contributions, domain entities, relations, and coreference structures in a machine-actionable form. However, most existing scientific information extraction resources mainly focus on entity categories, research topics, methods, tasks, and general semantic associations. They remain insufficient for characterizing fine-grained scientific logic relations, such as causal reasoning, parameter association, experimental verification, and innovation evaluation. Therefore, scientific text RE requires not only contextual semantic understanding, but also explicit modeling of domain-specific logic cues.

RE methods have evolved from rule-based and feature-engineered approaches to neural models and pre-trained language model frameworks. Early methods relied on manually designed patterns, syntactic features, or dependency structures. Neural models later improved RE performance through data-driven representation learning [19, 18, 11]. More recently, pre-trained language models such as BERT [4] and RoBERTa [12] have provided strong contextual semantic representations for RE, while SpanBERT further enhances span-level representation learning for entity and relation modeling [8]. Although these methods achieve promising performance on general-domain RE tasks, they mainly rely on implicit semantic association learning and lack explicit modeling of scientific logic, which limits their effectiveness in complex scientific texts.

To improve RE performance, end-to-end joint extraction methods have received increasing attention. CasRel formulates RE as a subject-first cascade tagging problem, first identifying subject entities and then predicting their corresponding relations and objects [17]. PRGC reduces redundant entity-pair matching through potential relation prediction and global correspondence learning [20]. REBEL reformulates RE as an end-to-end sequence generation task [5]. However, these methods mainly improve extraction through encoder or decoder design, while still relying on generic semantic representations and lacking explicit injection of scientific logic cues.

Scientific texts differ from general-domain texts in that entity relations are often embedded in argumentation structures rather than expressed through surface co-occurrence. Prior studies [10, 15] on scientific discourse have shown that scientific papers contain recognizable rhetorical and conceptual structures, such as background, methods, results, and conclusions. In such contexts, relations may be implied by causal reasoning, parameter association, experimental verification, or innovation evaluation. Without explicit logic-aware representation, models may confuse fine-grained scientific relations such as influence, support, verification, and evolution. Therefore, a key challenge in scientific text RE is how to inject domain-specific logic cues into contextual semantic representations. To address this challenge, FAST-RE introduces feature augmentation and logic-aware representation injection to enhance the extraction of implicit relations from complex scientific texts.

3 FAST-RE

As shown in Fig. 1, FAST-RE consists of six core modules.

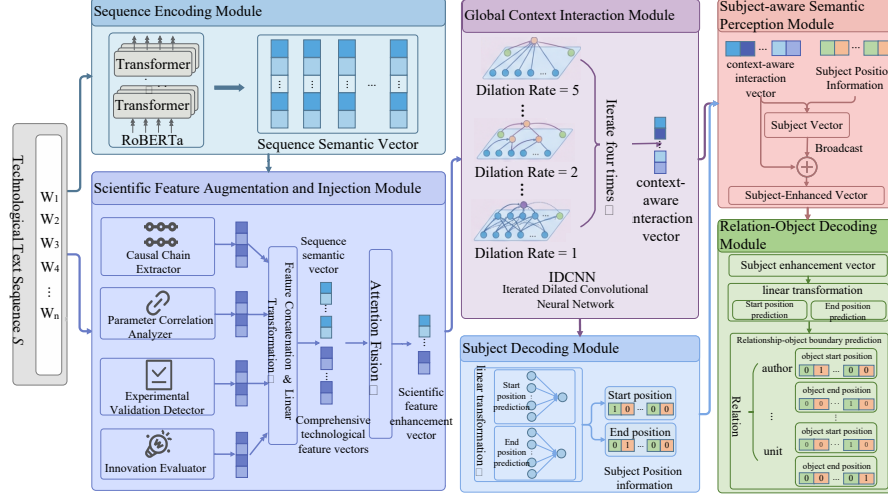


Fig. 1. The architecture of FAST-RE model

Given an input scientific sentence, the sequence encoding module first employs RoBERTa to generate contextual semantic representations. The scientific feature augmentation and injection module then captures explicit logic cues from four dimensions, namely causal reasoning, parameter association, experimental verification, and innovation evaluation, and dynamically injects them into the contextual representations to obtain feature-enhanced semantic vectors. Based on these enhanced representations, the global context interaction module introduces an iterated dilated convolutional network to model long-range contextual interactions and generate context-aware representations. Finally, FAST-RE adopts a subject-first decoding strategy for RE. The subject decoding module identifies subject entity boundaries, the subject-aware semantic perception module constructs subject-centered semantic representations, and the relation-object decoding module jointly predicts relation types and corresponding object entities.

3.1 Sequence Encoding Module

The sequence encoding module employs RoBERTa as the encoder to obtain contextual semantic representations of the input scientific text.

Given an input scientific text sequence $S = \{w_1, w_2, \dots, w_n\}$, where n represents the length of the input sequence, padding or truncation is applied to ensure a fixed input length. The sequence semantic vector H is obtained as:

$$H = \{h_1, h_2, \dots, h_n\} = \text{RoBERTa}(S) \quad (1)$$

where $H = \{h_1, h_2, \dots, h_n\}$, $h_i \in R^d$, h_i represents the contextual representation of the i -th token sequence, and d is the embedding dimension of RoBERTa. Unlike static word embeddings, RoBERTa dynamically models token semantics according to the surrounding context, which helps capture domain-specific terminology and complex contextual dependencies in texts.

3.2 Scientific Feature Augmentation and Injection Module

Scientific relations in technical texts are often not expressed by surface co-occurrence, but are implied by structured scientific logic. Relying only on generic contextual representations makes it difficult to explicitly capture such logic cues. To address this issue, FAST-RE introduces a scientific feature augmentation and injection module to extract multi-dimensional scientific logic features and inject them into token-level semantic representations.

Multi-dimensional Scientific Logic Modeling FAST-RE employs four scientific logic extractors to locate text fragments that carry domain-specific reasoning cues. Inspired by studies on scientific discourse [10, 15] and scientific knowledge organization [13, 1, 6], we model four types of scientific logic: causal reasoning, parameter association, experimental verification, and innovation evaluation. Let $m \in \{c, p, e, i\}$ denote the four dimensions, corresponding to causality, parameter, experiment, and innovation, respectively. The fragments activated by the m -th extractor are denoted as:

$$\mathcal{S}^{(m)} = \{S_1^{(m)}, S_2^{(m)}, \dots, S_{K_m}^{(m)}\} \quad (2)$$

where K_m is the number of activated fragments under dimension m .

Logic-aware Feature Extraction For the j -th fragment $S_j^{(m)}$, its start and end positions are represented as $s_j^{(m)}$ and $e_j^{(m)}$. Given the contextual representation $H = \{h_1, h_2, \dots, h_n\}$, the corresponding fragment representation is extracted as $H_j^m = [h_{s_j^{(m)}}^{(m)}, \dots, h_{e_j^{(m)}}^{(m)}]$. A dimension-specific BiGRU is then used to encode the contextual semantics within each fragment:

$$u_j^{(m)} = \text{BiGRU}^{(m)}(h_j) \quad (3)$$

where $u_j^{(m)}$ represents the semantic representation of the j -th fragment in the m -th scientific logic dimension.

Since different fragments contribute unequally to relation prediction, an attention-based aggregation mechanism is adopted to obtain the dimension-level logic feature:

$$\alpha_j^{(m)} = \frac{\exp\left(\mathbf{w}_j^{(m)} u_j^{(m)}\right)}{\sum_{k=1}^{K_m} \exp\left(\mathbf{w}_j^{(m)T} u_k^{(m)}\right)} \quad (4)$$

$$f^{(m)} = \sum_{j=1}^{K_m} \alpha_j^{(m)} u_j^{(m)} \quad (5)$$

where $\mathbf{w}^{(m)}$ is a trainable attention vector, $\alpha_j^{(m)}$ indicates the importance of the j -th fragment, and $f^{(m)}$ represents the logic-aware feature of dimension m .

Logic-aware Representation Injection The four-dimensional scientific logic features are concatenated to form the sentence-level scientific logic representation: $F = [f^{(c)}, f^{(p)}, f^{(e)}, f^{(i)}]$.

Then, a two-layer nonlinear transformation is applied to map the logic features into a unified representation space:

$$z = \sigma(W_1 F + b_1) \quad (6)$$

$$T = W_2 z + b_2 \quad (7)$$

where W_1 and W_2 are weight matrices, b_1 and b_2 are bias, and σ is the activation function. The vector T integrates causal reasoning, parameter association, experimental verification, and innovation evaluation cues into a unified scientific logic representation.

To inject scientific logic into token-level semantic representations, FAST-RE employs scaled dot-product attention. For the i -th token representation h_i , the query, key, and value are computed as:

$$T = [t^{(c)}; t^{(p)}; t^{(e)}; t^{(i)}] \in R^{4 \times d}, \quad (8)$$

$$q_i = W_Q h_i, \quad K = W_K T, \quad V = W_V T, \quad (9)$$

The logic-aware attention weight and supplementary feature vector are calculated as:

$$a_i = \text{softmax} \left(\frac{q_i K^\top}{\sqrt{d_k}} \right) \quad (10)$$

$$g_i = a_i v, \quad (11)$$

where d_k is the dimension of the key vector, and g_i denotes the scientific logic feature injected into the i -th token.

Finally, residual fusion and layer normalization are applied to obtain the feature-enhanced token representation:

$$\tilde{h}_i = \text{LayerNorm} (h_i + W_o[h_i; g_i]), \quad (12)$$

The output sequence of this module is $\tilde{H} = [\tilde{h}_1, \tilde{h}_2, \dots, \tilde{h}_n]$. Through this design, FAST-RE injects explicit scientific logic into contextual semantic representations, enabling token representations to capture both general semantics

and domain-specific reasoning cues. If no scientific logic fragment is activated, the injected feature is weakened, allowing $\tilde{H}$ to remain close to the original representation H , which improves robustness on conventional descriptive sentences.

3.3 Global Context Interaction Module

To further model long-range dependencies in scientific texts, FAST-RE introduces a global context interaction module based on the Iterated Dilated Convolutional Neural Network (IDCNN) [14]. Given the feature-enhanced sequence $\tilde{H} = \{\tilde{h}_1, \tilde{h}_2, \dots, \tilde{h}_n\}$, IDCNN enlarges the receptive field through stacked dilated convolutions while maintaining computational efficiency.

Inspired by Hybrid Dilated Convolution (HDC) [16], we adopt dilation rates $r \in \{1, 2, 5\}$ to reduce gridding effects and improve multi-range contextual coverage. For the l -th layer, the representation is computed as:

$$H^{(l)} = \sigma \left(\text{Conv}_{r_l}(H^{(l-1)}) \right), \quad (13)$$

where $H^{(0)} = \tilde{H}$, $\text{Conv}_{r_l}(\cdot)$ denotes a one-dimensional dilated convolution with dilation rate r_l , and σ is the ReLU activation function. The final context-aware representation is:

$$H_{ctx} = H^{(L)}. \quad (14)$$

The obtained H_{ctx} captures global contextual interactions while preserving logic-aware features, providing more informative representations for subject and relation-object decoding.

3.4 Subject Decoding Module

The subject decoding module takes the context-aware representation H_{ctx} as input and predicts subject entity boundaries in an end-to-end manner. Let $h_{ctx,i}$ denote the representation of the i -th token. Two independent linear classifiers are used to estimate the probabilities that each token serves as the start or end position of a subject entity:

$$P_i^{s,start} = \sigma(W_s^{start} h_{ctx,i} + b_s^{start}), \quad (15)$$

$$P_i^{s,end} = \sigma(W_s^{end} h_{ctx,i} + b_s^{end}), \quad (16)$$

where W_s^{start} and W_s^{end} are trainable weight matrices, b_s^{start} and b_s^{end} are bias, and $P_i^{s,start}$ and $P_i^{s,end}$ denote the probabilities that the i -th token is the start and end position of a subject entity, respectively.

During inference, tokens whose start or end probabilities exceed a predefined threshold τ are selected as candidate boundaries. A span (i, j) is identified as a subject entity if $P_i^{s,start} > \tau$, $P_j^{s,end} > \tau$, and $i \leq j$. All valid spans are collected to form the subject entity set:

$$\mathcal{S} = \{(i, j) \mid P_i^{s,start} > \tau, P_j^{s,end} > \tau, i \leq j\}. \quad (17)$$

3.5 Subject-aware Semantic Perception Module

After subject decoding, FAST-RE constructs a subject-aware semantic perception for each recognized subject. This module conditions token representations on the current subject, enabling subsequent relation-object decoding to focus on the relations associated with this specific subject.

For a recognized subject entity s , its start and end positions in H_{ctx} are denoted as p_{start} and p_{end} , respectively. We first convert p_{start} and p_{end} into one-hot vectors S_{head} and S_{tail} , whose lengths are equal to the maximum sequence length. The subject representation is then computed by averaging the contextual representations of its boundary tokens:

$$H_{sub} = \frac{\text{MatMul}(S_{head}, H_{ctx}) + \text{MatMul}(S_{tail}, H_{ctx})}{2}, \quad (18)$$

where $\text{MatMul}(\cdot)$ selects the corresponding boundary representations from H_{ctx} . In this way, H_{sub} captures the contextual semantics of the current subject.

To inject subject information into the sentence-level representation, H_{sub} is broadcast to each token position and fused with H_{ctx} :

$$H_p = H_{ctx} + \text{Broadcast}(H_{sub}), \quad (19)$$

where $\text{Broadcast}(\cdot)$ expands the subject representation to the same sequence length as H_{ctx} . The resulting representation H_p incorporates both global contextual information and subject-specific semantic bias, providing subject-aware features for relation-object decoding.

3.6 Relation-Object Decoding Module

The relation-object decoding module predicts relation types and object entities conditioned on each recognized subject. Since a subject may correspond to multiple relations and objects, FAST-RE performs relation-specific object boundary prediction based on the subject-aware representation H_p . For each relation type r , two linear classifiers predict the start and end probabilities of object entities:

$$P_i^{r,o,start} = \sigma(W_r^{o,start} h_{p,i} + b_r^{o,start}), \quad (20)$$

$$P_i^{r,o,end} = \sigma(W_r^{o,end} h_{p,i} + b_r^{o,end}), \quad (21)$$

where $h_{p,i}$ is the representation of the i -th token in H_p , and $W_r^{o,start}$, $W_r^{o,end}$, $b_r^{o,start}$, and $b_r^{o,end}$ are trainable parameters.

During inference, for each subject s and relation type r , candidate object spans are selected by:

$$\mathcal{O}_r = \{(i, j) \mid P_i^{r,o,start} > \tau, P_j^{r,o,end} > \tau, i \leq j\}. \quad (22)$$

Each valid span is mapped to an object entity o_k , and the final triples are obtained as:

$$\mathcal{T} = \{(s, r, o_k) \mid o_k \in \mathcal{O}_r, r \in \mathcal{R}\}, \quad (23)$$

where $\mathcal{R}$ is the predefined relation set. This subject-conditioned decoding strategy enables FAST-RE to extract multiple relation-object pairs for the same subject.

3.7 Loss Function

The total loss of FAST-RE consists of the subject decoding loss and the relation-object decoding loss:

$$\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{sub}} + \mathcal{L}_{\text{rel-obj}}. \quad (24)$$

The subject decoding loss is defined as:

$$\mathcal{L}_{\text{sub}} = - \sum_{t \in \{\text{start}, \text{end}\}} \sum_{i=1}^n [\gamma_1 y_i^{s,t} \log P_i^{s,t} + \gamma_0 (1 - y_i^{s,t}) \log(1 - P_i^{s,t})], \quad (25)$$

where n is the sequence length, $y_i^{s,t}$ and $P_i^{s,t}$ denote the ground-truth label and predicted probability for the i -th token being a subject boundary of type t , respectively. γ_1 and γ_0 weight positive and negative samples to alleviate class imbalance.

The relation-object decoding loss is computed over all relation types:

$$\mathcal{L}_{\text{rel-obj}} = - \sum_{r=1}^R \sum_{t \in \{\text{start}, \text{end}\}} \sum_{i=1}^n [\gamma_1 y_i^{r,o,t} \log P_i^{r,o,t} + \gamma_0 (1 - y_i^{r,o,t}) \log(1 - P_i^{r,o,t})], \quad (26)$$

where R is the number of relation types, and $y_i^{r,o,t}$ and $P_i^{r,o,t}$ denote the label and predicted probability for the i -th token being an object boundary of type t under relation r . This objective jointly optimizes subject extraction and relation-specific object decoding.

4 Experiments

4.1 Experimental Setup

Dataset Existing scientific-domain RE datasets mainly focus on entity categories, static attributes, and shallow semantic associations, with limited coverage of fine-grained scientific logic. To better evaluate RE models in complex scientific scenarios, we construct a fine-grained scientific text RE dataset based on a scientific knowledge ontology.

The corpus is collected from the abstracts and conclusions of computer science papers and processed through text cleaning, invalid sentence filtering, and long-sentence splitting. We select 10 high-frequency relation types covering method application, problem solving, parameter-performance interaction, experimental verification, and technical evolution. Preliminary annotations are generated by a large language model and then reviewed by three trained annotators in Label Studio, achieving a Fleiss' Kappa of 89.56% and a Kendall coefficient of 90.15%.

As shown in Table 1, the dataset contains 4,000 sentence-level samples and 6,120 annotated triples across 10 relation types, with a relatively balanced distribution. It includes 1,513 logically complex sentences involving causal reasoning, parameter association, and other explicit scientific logic, and 2,487 regular descriptive sentences. The dataset is split into training, validation, and test sets with a ratio of 8:1:1.

Table 1. Relation categories and data sample distribution.

Relation Type	Triples	Proportion (%)
Containment	665	10.87
Application	572	9.35
Solution	684	11.17
Output	585	9.56
Measurement	642	10.49
Influence	615	10.05
Verification	638	10.42
Support	575	9.40
Comparison	556	9.08
Evolution	588	9.61

Evaluation Metrics We evaluate all models using precision, recall, and $F1$ score, which are widely used metrics for relation extraction. A predicted triple is regarded as correct only when its subject, relation type, and object exactly match the ground-truth annotation. Precision, recall, and $F1$ score are defined as follows:

$$P = \frac{TP}{TP + FP}, \quad (27)$$

$$R = \frac{TP}{TP + FN}, \quad (28)$$

$$F1 = \frac{2PR}{P + R}, \quad (29)$$

where TP , FP , and FN denote the numbers of true positive, false positive, and false negative triples, respectively.

Implementation Details FAST-RE is implemented with PyTorch 1.13.1 and trained on an NVIDIA GeForce RTX 3090 GPU with 24GB memory. RoBERTa-wwm-ext (<https://huggingface.co/hfl/chinese-roberta-wwm-ext>) is used as the backbone encoder, with a hidden dimension of 768. The model is trained for 50 epochs using the Adafactor optimizer. The learning rate is set to 1×10^{-4} , and the dropout rate is 0.1. The training and testing batch sizes are 28 and 12, respectively. For the weighted binary cross-entropy loss, the positive and negative sample weights are set to $\gamma_1 = 1.0$ and $\gamma_0 = 0.3$.

4.2 Baseline Comparison

Baseline Models The baseline models are as follows:

NovelTagging [21] formulates joint entity and relation extraction as a sequence labeling task by designing a unified tagging scheme.

SpanBERT [8] enhances span-level representation learning through span-based pre-training objectives, making it suitable for entity-span-centered relation extraction.

PRGC [20] introduces potential relation prediction and global correspondence learning to reduce redundant entity-pair matching.

REBEL [5] reformulates RE as a sequence-to-sequence generation task and directly generates relation triples.

RA-CasRel [9] extends the CasRel framework with relation attention and relation coverage mechanisms to improve extraction under overlapping relation scenarios.

MICB [7] adopts a multi-layer indexing and cascaded binary classification framework to enhance nested entity recognition and complete triple extraction in a cascaded subject-relation-object manner.

Comparison Results As shown in Table 2, FAST-RE achieves the best precision and $F1$ score, reaching 81.053% and 78.062%, respectively.

Table 2. Performance Comparison of RE Results of Different Models

Model	Precision/%	Recall/%	$F1$ Score/%
NovelTagging	57.468	54.629	56.013
Spanbert	61.023	59.185	60.093
PRGC	69.829	66.726	68.242
REBEL	71.713	70.026	70.859
RA-CasRel	76.365	72.059	74.145
MICB	79.257	76.035	77.614
FAST-RE	81.053	75.284	78.062

Compared with the strongest baseline MICB, FAST-RE improves precision by 1.796% and $F1$ score by 0.448%. Although MICB obtains the highest recall, FAST-RE achieves better overall performance, indicating that logic-aware representation injection helps reduce incorrect relation predictions while maintaining competitive coverage. NovelTagging and SpanBERT perform relatively poorly, suggesting that simple sequence labeling or span-level semantic modeling is insufficient for complex scientific text RE. PRGC, REBEL, and RA-CasRel improve triple extraction through relation prediction, generative modeling, or subject-first decoding, but they still mainly rely on generic semantic representations and lack explicit scientific logic modeling. MICB achieves strong recall through cascaded binary classification and multi-layer indexing, but its broader candidate coverage may introduce additional false positives, resulting in lower precision than FAST-RE.

The advantage of FAST-RE mainly comes from its scientific feature augmentation and injection module. By modeling scientific logic cues such as causal reasoning, parameter association, experimental verification, and innovation evaluation, FAST-RE injects logic-aware features into contextual representations and guides the model to focus on relation clues consistent with scientific expression patterns. Meanwhile, the global context interaction module further enhances long-range semantic interaction over the augmented representations.

These results demonstrate that combining logic-aware representation injection with global context interaction improves the reliability and overall effectiveness of scientific text RE.

4.3 Ablation Study

We conduct component-level and subset-level ablation experiments to evaluate the contributions of Scientific Feature Augmentation and Injection (SFAI) module and the Global Context Interaction (GCI) module, and further analyze the effect of SFAI across different sentence expression patterns.

Component Ablation As shown in Table 3, removing either module leads to clear performance degradation, indicating that both SFAI and GCI contribute to the overall extraction capability of FAST-RE. When the GCI module is removed, the $F1$ score drops from 78.062% to 67.496%, with a decrease of 10.566%. This shows that global contextual interaction is important for capturing long-distance dependencies and supporting both subject boundary detection and relation-object decoding.

Table 3. Component ablation results.

Model	Precision (%)	Recall (%)	F1 (%)
FAST-RE w/o GCI	70.391	64.827	67.496
FAST-RE w/o SFAI	73.902	68.471	71.083
FAST-RE	81.053	75.284	78.062

When the SFAI module is removed, the $F1$ score decreases from 78.062% to 71.083%, with a drop of 6.979%. This result confirms that explicit scientific logic cues provide effective domain guidance for relation extraction. Without SFAI, the model can still rely on contextual encoding and global interaction, but its ability to recognize logic-sensitive relations is weakened, especially for expressions involving causal reasoning, parameter association, experimental verification, and innovation evaluation. Overall, GCI mainly enhances global semantic interaction, while SFAI introduces domain-specific logic guidance. Their combination enables FAST-RE to achieve more reliable extraction in complex scientific texts.

Fine-grained Subset Analysis To further examine SFAI, we divide the test set into a general descriptive subset and an explicit logic subset, accounting for 62.2% and 37.8%, respectively. The former contains method applications and factual statements, while the latter involves logic-intensive expressions such as causal reasoning and parameter association. The results are shown in Table 4.

As shown in Table 4, the contribution of SFAI varies across different text types. On the general descriptive subset, removing SFAI reduces the $F1$ score from 79.307% to 77.013%, with a decrease of 2.294%. This indicates that SFAI

Table 4. Fine-grained subset ablation results.

Model Variant	Subset Type	Precision (%)	Recall (%)	F1 (%)
FAST-RE	General Descriptive	82.306	76.518	79.307
FAST-RE w/o SFAI	General Descriptive	80.371	73.924	77.013
FAST-RE	Explicit Logic	79.358	73.192	76.150
FAST-RE w/o SFAI	Explicit Logic	63.509	59.784	61.590

provides auxiliary benefits for regular scientific descriptions by enriching semantic representations with scientific logic cues.

In contrast, on the explicit logic subset, removing SFAI causes a much larger $F1$ drop, from 76.150% to 61.590%, with a decrease of 14.560%. This result suggests that generic semantic representations and contextual interaction alone are insufficient for extracting implicit relations from logic-intensive scientific texts. By explicitly modeling causal reasoning, parameter association, experimental verification, and innovation evaluation, SFAI substantially improves the model’s ability to identify complex scientific relations.

These fine-grained results further demonstrate that SFAI does not merely increase model capacity, but provides targeted logic-aware enhancement for scientific texts with complex reasoning patterns. This verifies the necessity of explicit scientific logic modeling in FAST-RE.

5 Conclusion

In this paper, we propose FAST-RE, a Feature-Augmented Scientific Text RE model with logic-aware representation injection. FAST-RE explicitly captures four types of scientific logic, namely causal reasoning, parameter association, experimental verification, and innovation evaluation, and injects the resulting logic-aware features into semantic representations to enhance RE in scientific texts. The experimental results show that FAST-RE achieves the highest precision and $F1$ score compared to baselines. Further ablation studies confirm the effectiveness of the scientific feature augmentation and injection module, especially in handling complex scientific argumentation structures and implicit semantic relations. The performance gains are particularly evident for sentences involving complex reasoning patterns and long-distance dependencies. In future work, we will explore syntactic awareness and cross-domain transfer to improve generalization across complex and multi-domain scientific texts.

Acknowledgments This work is supported by the Natural Science Foundation of Jilin Province under grants 20240302083GX.

References

1. Auer, S., Kovtun, V., Prinz, M., Kasprzik, A., Stocker, M., Vidal, M.E.: Towards a knowledge graph for science. In: Proceedings of the 8th International Conference

- on Web Intelligence, Mining and Semantics. WIMS '18, Association for Computing Machinery, New York, NY, USA (2018). <https://doi.org/10.1145/3227609.3227689>, <https://doi.org/10.1145/3227609.3227689>
2. Bunescu, R., Mooney, R.J.: A shortest path dependency kernel for relation extraction. In: Proceedings of the Human Language Technology Conference and the Conference on Empirical Methods in Natural Language Processing. pp. 724–731. Association for Computational Linguistics, Vancouver, British Columbia, Canada (October 2005). <https://doi.org/10.3115/1220575.1220666>
 3. Dessì, D., Osborne, F., Reforgiato Recupero, D., Buscaldi, D., Motta, E., Sack, H.: Ai-kg: an automatically generated knowledge graph of artificial intelligence. In: International semantic web conference. pp. 127–143. Springer (2020)
 4. Devlin, J., Chang, M., Lee, K., Toutanova, K.: 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: Human Language Technologies, Volume 1 (Long and Short Papers). pp. 4171–4186. Association for Computational Linguistics, Minneapolis, Minnesota (2019). <https://doi.org/10.18653/v1/N19-1423>, <https://aclanthology.org/N19-1423/>
 5. Huguet Cabot, P.L., Navigli, R.: REBEL: Relation extraction by end-to-end language generation. In: Moens, M.F., Huang, X., Specia, L., Yih, S.W.t. (eds.) Findings of the Association for Computational Linguistics: EMNLP 2021. pp. 2370–2381. Association for Computational Linguistics, Punta Cana, Dominican Republic (Nov 2021). <https://doi.org/10.18653/v1/2021.findings-emnlp.204>, <https://aclanthology.org/2021.findings-emnlp.204/>
 6. Jaradeh, M.Y., Auer, S., Prinz, M., Kovtun, V., Kismihók, G., Stocker, M.: Open research knowledge graph: Towards machine actionability in scholarly communication. arXiv preprint arXiv:1901.10816 p. 42 (2019)
 7. Ji, W., Wen, K., Ding, L., Song, B.: A relation extraction method based on multi-layer index and cascading binary framework. In: International Conference on Advanced Data Mining and Applications (2024), <https://api.semanticscholar.org/CorpusID:275280561>
 8. Joshi, M., Chen, D., Liu, Y., Weld, D.S., Zettlemoyer, L., Levy, O.: SpanBERT: Improving pre-training by representing and predicting spans. Transactions of the Association for Computational Linguistics **8**, 64–77 (2020). https://doi.org/10.1162/tacl_a_00300, <https://aclanthology.org/2020.tacl-1.5/>
 9. Li, X., Li, Y., Yang, J., Liu, H., Hu, P.: A relation aware embedding mechanism for relation extraction. Applied Intelligence **52**(9), 10022–10031 (Jul 2022). <https://doi.org/10.1007/s10489-021-02699-3>, <https://doi.org/10.1007/s10489-021-02699-3>
 10. Liakata, M., Saha, S., Dobnik, S., Batchelor, C., Rebholz-Schuhmann, D.: Automatic recognition of conceptualization zones in scientific articles and two life science applications. Bioinformatics **28**(7), 991–1000 (04 2012). <https://doi.org/10.1093/bioinformatics/bts071>, <https://doi.org/10.1093/bioinformatics/bts071>
 11. Lin, Y., Shen, S., Liu, Z., Luan, H., Sun, M.: Neural relation extraction with selective attention over instances. In: Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers). pp. 2124–2133. Association for Computational Linguistics, Berlin, Germany (August 2016). <https://doi.org/10.18653/v1/P16-1200>, <https://aclanthology.org/P16-1200/>

12. Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., Stoyanov, V.: RoBERTa: A robustly optimized BERT pretraining approach. arXiv preprint arXiv:1907.11692 (2019), <https://arxiv.org/abs/1907.11692>
13. Luan, Y., He, L., Ostendorf, M., Hajishirzi, H.: Multi-task identification of entities, relations, and coreference for scientific knowledge graph construction. In: Riloff, E., Chiang, D., Hockenmaier, J., Tsujii, J. (eds.) Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing. pp. 3219–3232. Association for Computational Linguistics, Brussels, Belgium (Oct–Nov 2018). <https://doi.org/10.18653/v1/D18-1360>, <https://aclanthology.org/D18-1360/>
14. Strubell, E., Verga, P., Belanger, D., McCallum, A.: Fast and accurate entity recognition with iterated dilated convolutions. In: Palmer, M., Hwa, R., Riedel, S. (eds.) Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing. pp. 2670–2680. Association for Computational Linguistics, Copenhagen, Denmark (Sep 2017). <https://doi.org/10.18653/v1/D17-1283>, <https://aclanthology.org/D17-1283/>
15. Teufel, S., Moens, M.: Summarizing scientific articles: Experiments with relevance and rhetorical status. *Computational Linguistics* **28**(4), 409–445 (12 2002). <https://doi.org/10.1162/089120102762671936>, <https://doi.org/10.1162/089120102762671936>
16. Wang, P., Chen, P., Yuan, Y., Liu, D., Huang, Z., Hou, X., Cottrell, G.: Understanding convolution for semantic segmentation. In: 2018 IEEE Winter Conference on Applications of Computer Vision (WACV). pp. 1451–1460 (March 2018). <https://doi.org/10.1109/WACV.2018.00163>
17. Wei, Z., Su, J., Wang, Y., Tian, Y., Chang, Y.: A novel cascade binary tagging framework for relational triple extraction. In: Jurafsky, D., Chai, J., Schluter, N., Tetreault, J. (eds.) Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics. pp. 1476–1488. Association for Computational Linguistics, Online (Jul 2020). <https://doi.org/10.18653/v1/2020.acl-main.136>, <https://aclanthology.org/2020.acl-main.136/>
18. Zeng, D., Liu, K., Chen, Y., Zhao, J.: Distant supervision for relation extraction via piecewise convolutional neural networks. In: Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing. pp. 1753–1762. Association for Computational Linguistics, Lisbon, Portugal (September 2015). <https://doi.org/10.18653/v1/D15-1203>, <https://aclanthology.org/D15-1203/>
19. Zeng, D., Liu, K., Lai, S., Zhou, G., Zhao, J.: Relation classification via convolutional deep neural network. In: Proceedings of COLING 2014, the 25th International Conference on Computational Linguistics: Technical Papers. pp. 2335–2344. Dublin City University and Association for Computational Linguistics, Dublin, Ireland (August 2014), <https://aclanthology.org/C14-1220/>
20. Zheng, H., Wen, R., Chen, X., Yang, Y., Zhang, Y., Zhang, Z., Zhang, N., Qin, B., Ming, X., Zheng, Y.: PRGC: Potential relation and global correspondence based joint relational triple extraction. In: Zong, C., Xia, F., Li, W., Navigli, R. (eds.) Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers). pp. 6225–6235. Association for Computational Linguistics, Online (Aug 2021). <https://doi.org/10.18653/v1/2021.acl-long.486>, <https://aclanthology.org/2021.acl-long.486/>
21. Zheng, S., Wang, F., Bao, H., Hao, Y., Zhou, P., Xu, B.: Joint extraction of entities and relations based on a novel tagging scheme. In: Barzilay, R., Kan, M.Y. (eds.)

Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers). pp. 1227–1236. Association for Computational Linguistics, Vancouver, Canada (Jul 2017). <https://doi.org/10.18653/v1/P17-1113>, <https://aclanthology.org/P17-1113/>