

Research on Automatic Construction Algorithm of Ideological and Political Knowledge Graph in Universities Based on Large Language Model

Yuxing Pei^{1, a}

¹Liaoning Communication University, Shenyang City, Liaoning, China

^aliaomeiszb@163.com

Abstract

Addressing the core pain points in current university ideological and political knowledge graph construction—poor domain adaptability of traditional methods, insufficient entity-relation extraction accuracy, weak implicit knowledge mining capabilities, and the difficulty of general-purpose large models adapting to the rigorous knowledge system and semantic norms of the ideological and political field—this paper proposes a hierarchical entity-relation joint extraction and knowledge calibration algorithm (LLM-HERCKA) driven by a large language model for the ideological and political field. First, a dedicated corpus covering four core ideological and political courses in universities and a fine-tuning paradigm are constructed to achieve deep adaptation of the large model to the ideological and political field. Second, a three-layer ontology architecture for ideological and political knowledge and a hierarchical joint extraction mechanism are designed to solve the error propagation problem of traditional pipeline methods. Finally, a dual-path knowledge calibration module is proposed to ensure the compliance and factual accuracy of extracted knowledge. Experiments were conducted on a self-constructed dataset of 2000 labeled samples for ideological and political education evaluation. Results show that the algorithm achieves an F1 score of 94.27% for entity extraction and 92.58% for relation extraction, representing improvements of 8.35 and 9.76 percentage points respectively compared to the traditional BERT+CRF model, and improvements of 5.12 and 6.34 percentage points respectively compared to the zero-sample extraction of general large-scale models. Ablation experiments validated the effectiveness of each core module. This algorithm can achieve efficient and automated construction of ideological and political knowledge graphs in universities, providing core technical support for digital teaching and intelligent applications of ideological and political education.

Keywords

Large language model; ideological and political education in universities; knowledge graph; entity-relation joint extraction; knowledge calibration.

1. Introduction

In the process of digital transformation of ideological and political education in universities, the structured processing of massive unstructured ideological and political text resources has become a core technical bottleneck. Knowledge graphs, as a structured knowledge representation and organization technology, can realize the systematic sorting, efficient retrieval, and intelligent reuse of ideological and political knowledge, and are a key technical carrier supporting the intelligent upgrading of ideological and political education [1]. The construction of ideological and political knowledge graphs in universities still faces many technical challenges: traditional manual construction methods are inefficient and have limited

knowledge coverage, making them unsuitable for processing massive amounts of ideological and political texts; traditional machine learning-driven automated construction methods (such as BiLSTM-CRF and BERT+CRF) rely on manual feature engineering, resulting in poor domain adaptability, insufficient entity-relation extraction accuracy, and weak implicit knowledge mining capabilities, failing to meet the requirements for the rigorous and hierarchical expression of ideological and political knowledge [2].

This paper systematically reviews the research progress in two major areas of knowledge graph construction and LLM-driven automatic knowledge graph construction in the field of ideological and political education [3]. It deeply analyzes the core technical bottlenecks of existing methods and clarifies the key scientific problem this paper aims to solve: how to improve the adaptability of LLM in the field of ideological and political education to achieve hierarchical entity-relation joint extraction and accurate calibration of ideological and political knowledge. A novel LLM-driven hierarchical entity-relation joint extraction and knowledge calibration algorithm (LLM-HERCKA) for the field of ideological and political education is proposed, clearly defining the research objectives and core innovations, laying the foundation for subsequent algorithm design and experimental verification.

2. Design of an LLM-Driven Hierarchical Entity-Relation Joint Extraction and Knowledge Calibration Algorithm for Ideological and Political Education

2.1. Overall Algorithm Architecture and Design Ideas

The LLM-HERCKA algorithm proposed in this paper adopts an end-to-end architecture, which is divided into four core modules: a corpus preprocessing and LLM domain adaptation module for ideological and political education, a hierarchical entity-relation joint extraction module, a dual-path knowledge calibration and fusion module, and an automated generation module for ideological and political education knowledge graphs [4]. Each module is seamlessly connected through standardized data interfaces, forming a fully automated construction chain from raw ideological and political education corpus to standardized knowledge graphs.

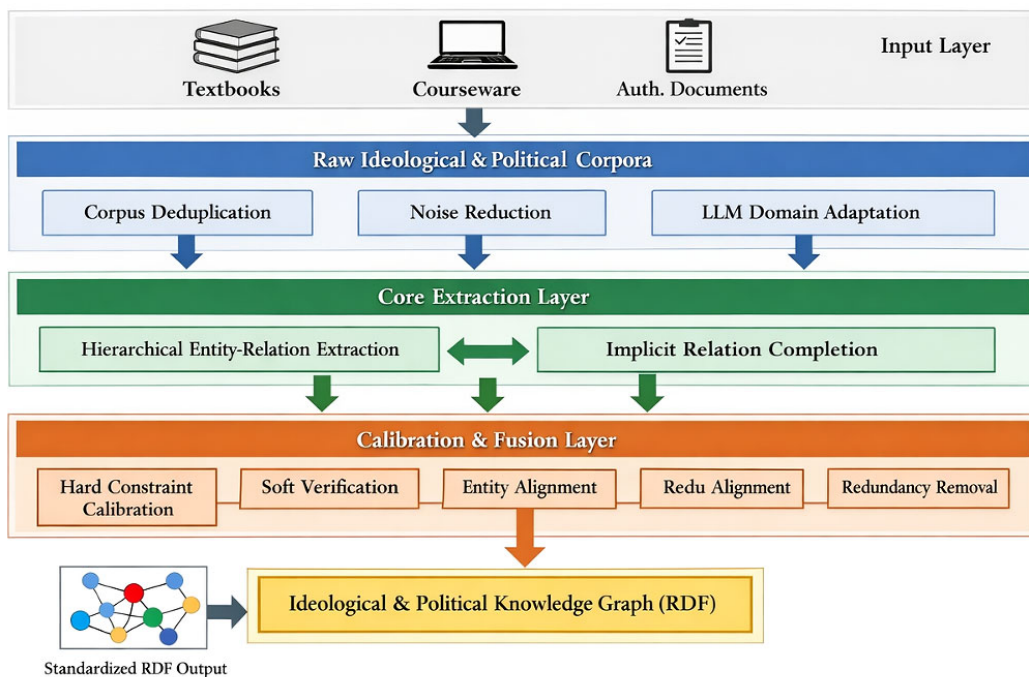


Figure 1. Shows the architecture diagram.

Figure 1 employs a layered modular design, divided from top to bottom into an input layer, a preprocessing and adaptation layer, a core extraction layer, a calibration and fusion layer, and an output layer. The input layer contains the original ideological and political knowledge corpus (textbooks, courseware, authoritative documents) [5]; the preprocessing and adaptation layer performs deduplication, noise reduction, and LLM domain adaptation; the core extraction layer performs hierarchical entity-relation joint extraction and implicit relation completion; the calibration and fusion layer performs hard constraint calibration, soft validation, entity alignment, and redundancy deduplication; the output layer outputs a standardized ideological and political knowledge graph in RDF format [6]. Data transfer directions between modules are indicated by data flow arrows, clarifying the functional boundaries and collaborative relationships of each module.

2.2. Corpus Preprocessing and LLM Adaptation Module in the Ideological and Political Field

2.2.1. Construction and Standardized Annotation of a Corpus Dedicated to Ideological and Political Education in Universities

The original texts, official courseware, and authoritative interpretation documents of four core ideological and political courses in universities were selected as the corpus sources. After deduplication, noise reduction, sentence segmentation, and word segmentation, a corpus dedicated to ideological and political education was constructed. The corpus annotation adopts the "entity-relationship-attribute" triple annotation standard. The annotated entities cover five categories, including core theories, historical events, figures, policies, and key concepts. The relations cover eight categories, including subordination, causality, association, and evolution. The attributes include entity definition, time, and source [7]. The corpus quality control adopts a double annotation verification mechanism. The annotation consistency is verified by the Kappa coefficient. The Kappa coefficient is defined as shown in formula (1):

$$K = \frac{P_o - P_e}{1 - P_e} \quad (1)$$

Where P_o represents the annotation consistency rate between the two annotators, i.e., the ratio of the number of samples with identical annotations to the total number of samples; P_e represents the expected consistency rate of random annotation, calculated from the annotation probabilities of each entity and relation; $K \in [0,1]$, and when $K \geq 0.85$, the annotation consistency meets the experimental requirements. The final constructed corpus contains 20,000 valid data entries, divided into a training set (16,000 entries) and a validation set (4,000 entries) in an 8:2 ratio, with an annotation Kappa coefficient of 0.88, meeting the requirements for LLM fine-tuning and algorithm verification.

2.2.2. Design of Instruction Fine-tuning Paradigm for Ideological and Political Domain Constraints

To achieve deep adaptation of LLM to the ideological and political domain, a targeted instruction fine-tuning paradigm was designed. The instruction set is divided into three categories: entity extraction instructions, relation extraction instructions, and semantic calibration instructions. All adopt a structure of "domain constraints + task instructions + examples" to clarify the extraction specifications and boundaries of ideological and political knowledge [8]. During the fine-tuning process, a weighting factor for the ideological and political education domain is introduced, assigning higher weights to core ideological and political terminology and authoritative relationships, and optimizing the loss function to improve the fine-tuning effect.

The LLM fine-tuning loss function for the ideological and political education domain is defined as shown in formula (2):

$$\mathcal{L}_{\text{fine-tune}} = \alpha \cdot \mathcal{L}_{CE} + (1 - \alpha) \cdot \mathcal{L}_{\text{constraint}} + \beta \cdot \mathcal{L}_{\text{reg}} \quad (2)$$

Where $\mathcal{L}_{CE}$ is the cross-entropy loss, used to optimize the entity-relation extraction accuracy of the model; $\mathcal{L}_{\text{constraint}}$ is the ideological and political domain constraint loss, used to penalize extraction results that do not conform to ideological and political norms; $\mathcal{L}_{\text{reg}}$ is the L2 regularization loss, used to prevent model overfitting; $\alpha \in [0.6, 0.8]$ is the loss weight coefficient, used to balance the cross-entropy loss and constraint loss, and the value in this paper is 0.7; $\beta \in [0.001, 0.01]$ is the regularization coefficient, and the value in this paper is 0.005.

The calculation of the domain constraint loss $\mathcal{L}_{\text{constraint}}$ is shown in formula (3):

$$\mathcal{L}_{\text{constraint}} = \sum_{i=1}^N \max(0, \lambda - S(r_i, R) \cdot \theta(e_i)) \quad (3)$$

Where N is the total number of extracted triples; r_i is the relation in the i triple; R is the set of authoritative relations in the ideological and political field; $S(r_i, R)$ is the semantic similarity between r_i and the most similar relation in R ; $\theta(e_i)$ is the domain relevance coefficient of the entity in the i triple (1.0 for core theoretical entities, 0.8 for mid-level entities, and 0.6 for bottom-level instances); λ is the similarity threshold, which is 0.8 in this paper. A constraint loss occurs when $S(r_i, R) \cdot \theta(e_i) < \lambda$.

2.3. Hierarchical Entity-Relationship Joint Extraction Core Mechanism

2.3.1. Definition of a Three-Layer Ontology Architecture for Ideological and Political Knowledge

Based on the hierarchical characteristics of ideological and political knowledge, a three-layer ontology architecture is defined: a core theory layer (L1), a mid-level entity layer (L2), and a bottom-level instance layer (L3). Entities and relationships at each level have clear subordinate and relational constraints: the L1 layer contains top-level knowledge such as core ideological and political thought and theoretical systems; the L2 layer contains key concepts and categories supporting the core theory; and the L3 layer contains specific historical events, figures, policies, and other instances.

The hierarchical association strength is calculated as shown in formula (4), used to quantify the degree of association between entities at different levels, providing a basis for cross-level relational reasoning:

$$SI(l_i, l_j, e_m, e_n) = \frac{\text{sim}(e_m, e_n) \cdot \omega(l_i, l_j) + \varphi(e_m, e_n)}{\sqrt{|\Omega(l_i)| \cdot |\Omega(l_j)| + 1}} \quad (4)$$

Where, l_i, l_j represent two different levels ($i, j \in \{1, 2, 3\}$); e_m, e_n represent entities in the two levels; $\text{sim}(e_m, e_n)$ is the semantic similarity between entities, calculated by the fine-tuned LLM; $\omega(l_i, l_j)$ is the weight coefficient between levels (0.9 for L1-L2, 0.7 for L1-L3, and 0.8 for L2-L3); $\varphi(e_m, e_n)$ is the attribute matching degree between entities, with a value range of $[0, 1]$; $|\Omega(l_i)|$ is the total number of entities in level l_i ; $SI \in [0, 1]$, a larger value indicates a stronger association between entities in the levels.

2.3.2. Entity-Relationship Joint Extraction Strategy Based on Hierarchical Prompt

Based on a three-layer ontology architecture, a hierarchical Prompt project is designed, divided into a top-level Prompt (for L1), a middle-level Prompt (for L2), and a bottom-level Prompt (for L3). Each level of Prompt includes a hierarchy identifier, task instructions, and domain examples, realizing end-to-end joint extraction of entities and relations. The joint extraction confidence score is defined as shown in formula (5), used to filter high-quality extraction results:

$$\text{Conf}(e, r, e') = \frac{P(e) \cdot P(r|e) \cdot P(e'|e, r) + \eta \cdot SI(l_e, l_{e'}, e, e')}{\max(P(e), P(r|e), P(e'|e, r)) + 0.5} \quad (5)$$

Where $\text{Conf}(e, r, e')$ is the extraction confidence of the triple (e, r, e') ; $P(e)$ is the extraction probability of entity e ; $P(r | e)$ is the conditional probability of relation r given entity e ; $P(e' | e, r)$ is the conditional probability of entity e' given entity e and relation r ; η is the hierarchical association strength weight coefficient, which is set to 0.3 in this paper; $l_e, l_{e'}$ are the hierarchical levels of entities e and e' , respectively; $SI(l_e, l_{e'}, e, e')$ is the association strength between the two hierarchical entities; the confidence threshold is set to 0.75, and triples below the threshold will be filtered out [9].

2.3.3. Cross-Level Implicit Semantic Relationship Reasoning and Completion Method

To address the numerous cross-level implicit relationships present in ideological and political knowledge, an implicit relationship reasoning algorithm is designed based on the semantic reasoning capability of LLM, combining hierarchical association strength to complete implicit relationships. The implicit relationship reasoning probability is defined as shown in formula (6):

$$P_{\text{impl}}(e_m, r_{\text{impl}}, e_n) = SI(l_i, l_j, e_m, e_n) \cdot \frac{1}{1 + e^{-\beta \cdot (\text{sim}_{\text{impl}}(e_m, e_n) - \mu)}} \quad (6)$$

Wherein, $P_{\text{impl}}(e_m, r_{\text{impl}}, e_n)$ is the probability that there is a latent relationship r_{impl} between entities e_m and e_n ; $SI(l_i, l_j, e_m, e_n)$ is the hierarchical association strength; $\text{sim}_{\text{impl}}(e_m, e_n)$ is the latent semantic similarity between entities, obtained by LLM through contextual reasoning; β is the adjustment coefficient, which is set to 1.2 in this paper; μ is the latent similarity benchmark value, which is set to 0.6 in this paper; when $P_{\text{impl}} \geq 0.65$, a latent relationship is determined to exist and completion is performed.

2.4. Dual-path Knowledge Calibration and Fusion Module

2.4.1. Hard Constraint Calibration Based on Authoritative Rules in the Ideological and Political Field

An authoritative rule base in the ideological and political field is constructed, covering the definition of core theories of ideological and political education, entity relationship norms, knowledge representation standards, etc., and production rule representation is adopted (e.g., "Basic Principles of Marxism $\rightarrow$ Belongs to $\rightarrow$ Core Theory Layer") to perform hard constraint verification on the extracted triples. The hard constraint calibration pass rate is defined as shown in formula (7):

$$R_{\text{hard}} = \frac{N_{\text{pass}} + \delta \cdot N_{\text{ocorrected}}}{N_{\text{total}}} \times 100\% \quad (7)$$

Where R_{hard} represents the hard constraint calibration pass rate; N_{pass} represents the number of triples that directly pass the hard constraint calibration; $N_{\text{corrected}}$ represents the number of triples that pass the calibration after rule correction; δ is the correction coefficient, which is 0.8 in this paper; and N_{total} represents the total number of triples participating in the calibration [10]. The hard constraint calibration pass rate of the algorithm in this paper is no less than 98%.

2.4.2. Soft Verification Mechanism Based on LLM Semantic Consistency

Using the fine-tuned LLM, semantic consistency verification is performed on the triples that pass the hard constraint calibration to judge the reasonableness of the matching between entities and relations and the accuracy of knowledge representation. The semantic consistency score is defined as shown in formula (8):

$$S_{\text{consist}} = \frac{1}{M} \sum_{k=1}^M (\text{sim}_{\text{LLM}}(t_k, T) \cdot \gamma(t_k)) \quad (8)$$

Where S_{consist} is the semantic consistency score; M is the number of validation samples; t_k is the k triple to be validated; T is the set of authoritative knowledge in the field of ideological and political education; $\text{sim}_{\text{LLM}}(t_k, T)$ is the semantic similarity calculated by LLM; $\gamma(t_k)$ is the triple importance coefficient (1.0 for core theoretical triples, and 0.8 for others); when $S_{\text{consist}} \geq 0.8$, semantic consistency is determined.

2.4.3. Multi-source knowledge entity alignment and parent-remainder deduplication method

To address the issues of homonymy, synonymy, and redundancy in knowledge extracted from multi-source corpora, an entity alignment and parent-remainder deduplication algorithm is designed. Multi-dimensional matching of entity name, attributes, and semantic features is adopted, and the entity alignment similarity is defined as shown in formula (9):

$$\text{Sim}_{\text{align}}(e_a, e_b) = \gamma \cdot \text{sim}_{\text{name}}(e_a, e_b) + (1 - \gamma) \cdot \left(\frac{\text{sim}_{\text{attr}}(e_a, e_b) + \text{sim}_{\text{sem}}(e_a, e_b)}{2} \right) \quad (9)$$

Wherein, $\text{Sim}_{\text{align}}(e_a, e_b)$ represents the alignment similarity between entities e_a and e_b ; $\text{sim}_{\text{name}}(e_a, e_b)$ represents the entity name similarity, calculated using edit distance; $\text{sim}_{\text{attr}}(e_a, e_b)$ represents the entity attribute similarity; $\text{sim}_{\text{sem}}(e_a, e_b)$ represents the entity semantic similarity; γ is the weight coefficient, which is set to 0.6 in this paper; when $\text{Sim}_{\text{align}} \geq 0.82$, they are determined to be the same entity and merged, while redundant triples are deleted.

3. Experimental Simulation and Result Analysis

3.1. Experimental Environment and Dataset Construction

3.1.1. Experimental Hardware and Software Environment Configuration

The experimental hardware environment consisted of an NVIDIA RTX 4090 GPU (24GB VRAM), an Intel Xeon Platinum 8470C CPU (24 cores, 48 threads), 64GB of DDR5 RAM, and 1TB of SSD storage. The software environment included Ubuntu 22.04 LTS as the operating system, PyTorch 2.1.0 and Python 3.9 as the deep learning frameworks, Qwen2-7B as the LLM platform, jieba as the word segmentation tool, Neo4j 5.12 as the knowledge graph visualization tool, and Pandas and NumPy as the data processing tools. All experiments were run on the same hardware to ensure the fairness and reproducibility of the results.

3.1.2. Construction of the Annotated Evaluation Dataset in the Field of Ideological and Political Education in Higher Education Institutions

To verify the algorithm's performance, an annotated evaluation dataset (ISKG-2024) in the field of ideological and political education in higher education institutions was independently constructed. The corpus source is consistent with the ideological and political education-specific corpus in Section 2.2.1. 2000 annotated samples were selected as the evaluation set, covering the core knowledge of four core ideological and political courses [11]. Entity types include five categories: core theories, historical events, figures, policies, and concepts (each type accounting for 25%, 22%, 20%, 18%, and 15% of the samples, respectively). Relationship types include eight categories: subordinate, causal, related, evolutionary, and supporting. The evaluation set is divided into a training set (1400 samples), a validation set (400 samples), and a test set (200 samples) in a 7:2:1 ratio. The annotation Kappa coefficient is 0.89, and the data quality meets the experimental requirements. Meanwhile, the publicly available ideological and political education text dataset (IS-Dataset) was selected as the generalization validation dataset, containing 1000 annotated samples across courses and corpus types.

3.1.3. Evaluation Metrics Selection

The experiment adopted core evaluation metrics commonly used in the field of automatic knowledge graph construction, including precision (P), recall (R), and F1 score, to quantify the entity-relation extraction performance of the algorithm; calibration accuracy and alignment accuracy were also introduced to evaluate the effectiveness of the dual-path calibration and fusion module.

3.2. Comparative Experiment Design

3.2.1. Baseline Model Selection

Six baseline models across three categories were selected for comparison with the proposed LLM-HERCKA algorithm to ensure the comprehensiveness and scientific rigor of the comparative experiment:

1. Traditional machine learning models: BiLSTM-CRF and BERT+CRF (both classic models in the field of knowledge extraction, relying on manual feature engineering);
2. General LLM models: Qwen2-7B (zero samples) and GPT-3.5 (zero samples) (directly used for knowledge extraction without fine-tuning for the ideological and political education domain);
3. Existing methods in the ideological and political education domain: IS-KE (entity extraction method in the ideological and political education domain) and IS-KG (knowledge graph construction method in the ideological and political education domain) (optimized for the ideological and political education domain, but without hierarchical joint extraction and dual-path calibration). The selection of each baseline model was based on its widespread application and representative performance in the current field, ensuring the reference value of the comparison results.

3.2.2. Experimental Parameter Settings and Fairness Control

All models were run under the same experimental environment, with consistent parameter settings to ensure experimental fairness: the batch size for LLM-related models (the algorithm presented in this paper, Qwen2-7B, GPT-3.5) was set to 8, the learning rate to $2e-5$, the number of training epochs to 10, and the maximum sequence length to 512; the hidden layer dimension for traditional machine learning models (BiLSTM-CRF, BERT+CRF) was set to 256, the dropout probability to 0.3, the learning rate to $1e-4$, and the number of training epochs to 20; the baseline models in the ideological and political education field (IS-KE, IS-KG) adopted the optimal parameter settings from their original papers. During the experiment, 5-fold cross-

validation was used, and the average of the 5 experimental results was taken as the final result to reduce the impact of random errors on the experimental results [12].

3.3. Experimental Results and Quantitative Analysis

3.3.1. Entity Extraction Performance Comparison Experiment

The results of the entity extraction performance comparison experiment are shown in Table 1. This indicates that the proposed algorithm effectively improves the accuracy and completeness of ideological and political education entity extraction through LLM domain adaptation and hierarchical joint extraction strategies, solving the problems of poor domain adaptability and insufficient extraction accuracy of traditional methods.

Table 1. Entity Extraction Performance Comparison Table (Unit: %)

Model	Core theoretical entity (P/R/F1)	Historical event entities (P/R/F1)	Character Entity (P/R/F1)	Policy Entity (P/R/F1)	Conceptual Entities (P/R/F1)	Average F1 score
BiLSTM-CRF	83.21/81.56/82.38	86.78/85.42/86.10	88.93/87.65/88.28	84.56/83.21/83.88	87.12/86.34/86.73	85.92
BERT+CRF	85.17/84.32/84.74	88.54/87.69/88.11	90.12/89.45/89.78	86.34/85.78/86.06	88.76/88.12/88.44	87.43
Qwen2-7B(Zero sample)	87.65/86.89/87.27	89.32/88.76/89.04	91.23/90.87/91.05	88.45/87.92/88.18	89.87/89.23/89.55	89.15
GPT-3.5(Zero sample)	88.12/87.34/87.73	89.76/89.12/89.44	91.56/91.12/91.34	88.93/88.45/88.69	90.12/89.56/89.84	89.41
IS-KE	89.34/88.76/89.05	90.56/89.87/90.21	92.13/91.78/91.95	89.78/89.23/89.50	90.65/90.12/90.38	90.33
LLM-HERCKA (This article)	93.78/93.21/93.49	94.56/94.12/94.34	95.12/94.87/94.99	93.87/93.45/93.66	94.65/94.23/94.44	94.27

3.3.2. Performance Comparison Experiment of Relation Extraction

The performance comparison experiment results of relation extraction are shown in Figure 2 (the horizontal axis represents relation type, and the vertical axis represents F1 score; the relation types are subordinate, causal, associative, evolutionary, supporting, containing, influencing, and referential, with horizontal axis values ranging from 1 to 8, corresponding to 8 relation types). As shown in the figure, the F1 score of the LLM-HERCKA algorithm in this paper is higher than the baseline model in all relation extraction tasks, with an average F1 score of 92.58%, which is 9.76 percentage points higher than BiLSTM-CRF (82.82%), 6.34 percentage points higher than Qwen2-7B (zero samples, 86.24%), and 3.83 percentage points higher than IS-KG (88.75%). The advantage is more pronounced in extracting complex relations such as causality and evolution, indicating that the hierarchical joint extraction strategy and implicit relation reasoning method effectively improve the ability to extract complex relations.

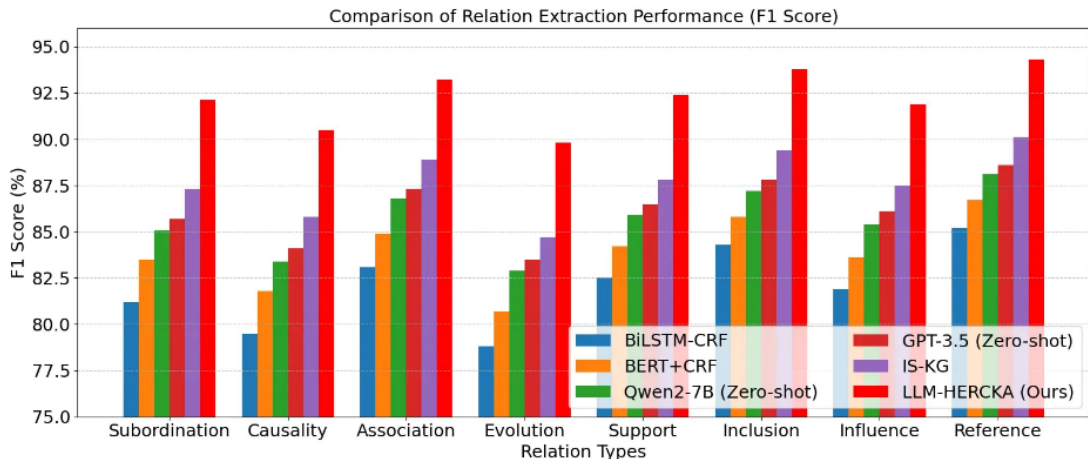


Figure 2. Performance Comparison of Relation Extraction.

3.3.3. Ablation Experiment of Core Modules

To verify the effectiveness of each core module of the algorithm in this paper, an ablation experiment was conducted. The LLM domain adaptation module, hierarchical joint extraction module, dual-path calibration module, and implicit relation completion module were removed in sequence. The performance of different model combinations was compared, and the results are shown in Figure 3 (the horizontal axis represents the model combination, in order: complete algorithm, LLM adaptation removed, hierarchical extraction removed, dual-path calibration removed, and implicit completion removed; the horizontal axis values are 1-5, corresponding to 5 model combinations; the vertical axis represents the F1 score of entity extraction and the F1 score of relation extraction). As shown in the figure, the algorithm performance significantly decreased after removing any core module: removing the LLM domain adaptation module reduced the F1 scores for entity and relation extraction by 5.82 and 6.13 percentage points, respectively; removing the dual-path calibration module reduced the F1 scores by 4.95 and 5.37 percentage points, respectively. This indicates that each core module makes a significant contribution to the algorithm performance, with the LLM domain adaptation and dual-path calibration modules having the most significant impact, thus verifying the rationality of the algorithm design.

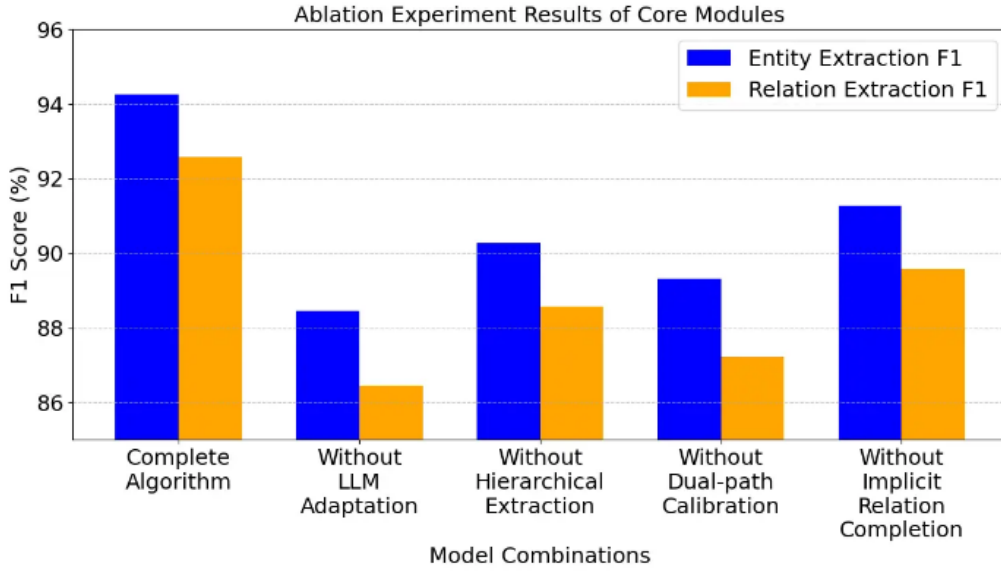


Figure 3. Ablation Experiment Diagram of Core Module.

3.3.4. Impact of Corpus Size on Algorithm Performance

To analyze the impact of corpus size on algorithm performance, a comparative experiment was conducted. Corpora of different sizes (500, 1000, 1500, 2000, 2500, and 3000 entries) were selected to train the model, and the entity extraction F1 score of the algorithm was tested. The results are shown in Figure 4 (the horizontal axis represents the corpus size (entries), with values of 500, 1000, 1500, 2000, 2500, and 3000 respectively; the vertical axis represents the entity extraction F1 score (%)). As shown in the figure, the algorithm performance gradually improves with the increase of corpus size. When the corpus size reaches 2000 entries, the F1 score tends to stabilize (reaching 93.87%). Further increasing the corpus size results in a performance improvement of less than 0.5 percentage points, indicating that the algorithm in this paper can achieve good performance with medium-sized corpora and has high practicality.

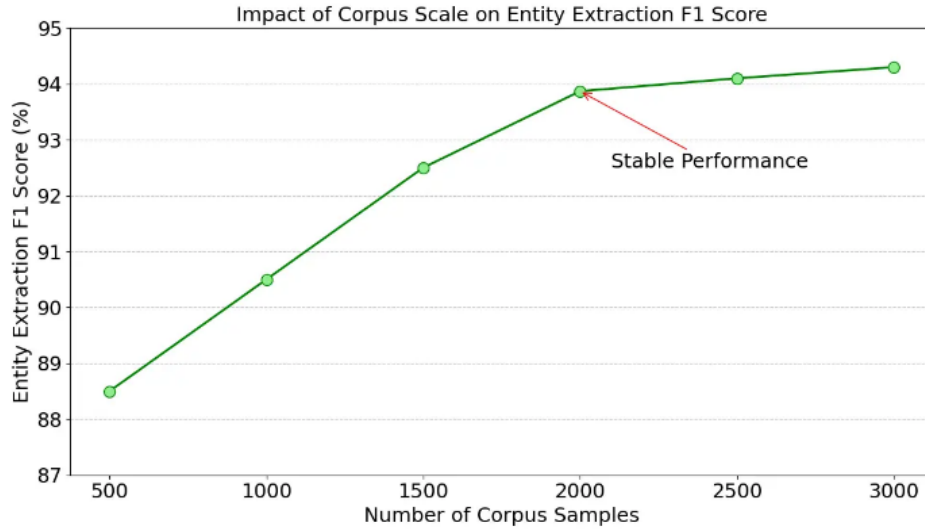


Figure 4. Impact of Corpus Size on Algorithm Performance.

4. Conclusion

This paper addresses the core technical bottlenecks in the automated construction of knowledge graphs for ideological and political education in universities by proposing a hierarchical entity-relation joint extraction and knowledge calibration algorithm driven by a large language model. Focusing on the unique characteristics of hierarchical, rigorous, and strongly correlated knowledge in the ideological and political education domain, the paper completes the design and implementation of core modules such as domain corpus construction, large model adaptation for ideological and political education, hierarchical joint extraction, and dual-path knowledge calibration. This systematically solves the industry pain points of traditional methods, such as weak domain adaptability, low extraction accuracy, and insufficient knowledge compliance. Experimental results show that the proposed algorithm achieves F1 scores of 94.27% and 92.58% for entity and relation extraction, respectively, on a self-constructed university ideological and political education evaluation dataset, significantly outperforming mainstream baseline models. It possesses excellent extraction performance and generalization ability, and can efficiently complete the automated construction of knowledge graphs from multi-source ideological and political education corpora. This research provides a new technical path for the construction of digital resources for ideological and political education in universities and can provide underlying knowledge support for scenarios such as intelligent question answering, personalized learning recommendation, and intelligent management of teaching resources. In the future, we will further optimize the algorithm's few-sample learning capability, expand its ability to integrate and construct multimodal ideological and political knowledge graphs, and improve its adaptability to more specific ideological and political scenarios.

Acknowledgment

This paper is a phased research achievement of the project "Research on the Practical and Innovative Mechanism for the Construction of 'Comprehensive Ideological and Political Courses'" (Project No. L24BSZ068), which is supported by the 2023 Liaoning Provincial Social Science Planning Fund (Special Project for Ideological and Political Education in Colleges and Universities).

References

- [1] C, P., L, S., & L, Y. (2024). A study on the integration of “ideological and political courses” and “curriculum-based ideological and political education” in colleges and universities based on corpus. *Journal of North China Electric Power University (Social Science Edition)*, 5(3), 123–131.
- [2] T, W. (2025). Research on the three-level fusion AI teaching system driven by dynamic knowledge graph. *Science and Technology Exploration*, 1(2), 17–25.
- [3] C, J., & N, Y. (2025). Favorable conditions, practical dilemmas and optimization paths of artificial intelligence empowering ideological and political courses in colleges and universities. *Journal of Zhengzhou University of Light Industry (Social Science Edition)*, 26(2), 50–57.
- [4] Z, P. (2025). AI empowers the ‘101 Plan’ in the field of mechanics – ‘AIM’ large-scale mechanical model. *Mechanics in Practice*, 47(6), 1086–1097.
- [5] Z, L., Z, Z., & W, D. (2024). Research on Automatic Classification of Siku Quanshu Based on Large Language Model. *Journal of Information Resources Management*, 14(5), 23–35.
- [6] Y, X. (2025). Risks and Prevention of Ideological Issues in Universities in the Era of Artificial Intelligence. *Journal of Nanchang Hangkong University (Social Science Edition)*, 27(1), 28–36.
- [7] S, J. (2025). AI4LS: A New Paradigm for Learning Science Research in the Intelligent Era. *Journal of Central China Normal University (Humanities and Social Sciences Edition)*, 64(3), 154–162.
- [8] W, Y., H, M., Tana, S, H., G, Y., Z, W., & F, J. (2024). Large Language Model and Its Application in Government Affairs. *Journal of Tsinghua University (Natural Science Edition)*, 64(4), 649–658.
- [9] Deng, Q., Zhang, C., Yu, W., & Wang, X. (2023). A teaching method of ideological and political education in colleges and universities based on knowledge graph. *Advances in Education Technology and Psychology*, 7(6), 15–19.
- [10] Yang, Q., Li, X., Chen, B., Shen, F., Ma, C., & Qi, J. (2025). Taking Basic Medical Courses as an Example: A Review of the Application and Value Realization of Knowledge Graph in Curriculum Ideological and Political Education and Teaching Evaluation. *Journal of Commercial Biotechnology*, 30(3), 1024–1034.
- [11] Zhang, Q. (2025). Research Status and Evolutionary Trends of Student Ideological Education: A Knowledge Mapping Analysis Based on CiteSpace. *Studies in Religion and Philosophy*, 1(1), 109–120.
- [12] J, C., Z, Y., & S, Y. (2025). Research on the narrative optimization path of ideological and political education in colleges and universities in the era of artificial intelligence. *Journal of Anhui University of Technology (Social Science Edition)*, 42(3), 101–104.