

Research on the Construction and Appreciation of a Knowledge Graph of Chinese Classical Literature and Culture for International Chinese Language Education

Yu Song^{1,*}

¹ School of International Education, Ningxia Medical University, Yinchuan, Ningxia 750004, China

* Correspondence author: catherine-06104@163.com

Abstract: As an effective means of gaining a deeper understanding of literary works, this paper proposes the construction of a knowledge graph to assist students in international Chinese language education with the interpretation of classical Chinese literature. The paper provides a comprehensive overview of the knowledge graph construction process and designs a BERT-BiLSTM-CRF model for entity recognition, using labeled sequences as the model's annotation results. Additionally, a Bs-Spert model is designed for entity relationship extraction, comprising a BERT pre-training module, a clustered search module, a span filtering module, and a relationship classification module. Finally, a knowledge Q&A model is built based on the knowledge graph embeddings to implement the application of the literary graph. The Chinese classical literature knowledge graph includes 12 entity types. The recognition accuracy for different entities is relatively high, with an overall recognition accuracy of 92.65%, a recall rate of 93.63%, and an F1 score of 90.47%. The model demonstrates relatively high performance across various precision metrics for knowledge extraction and exhibits good robustness. In knowledge-based question-answering, the model successfully extracts relevant information, achieving a match rate of 58.81%. The established knowledge graph promotes the inheritance and development of Chinese classical culture in international Chinese language education.

Keywords: knowledge graph; BERT-BiLSTM-CRF model; Bs-Spert model; knowledge-based question-answering; entity recognition

1. Introduction

International Chinese language education serves as the primary channel for the global dissemination of the Chinese language. The discipline possesses a strong sense of mission and inherent advantages in promoting Chinese culture, playing an irreplaceable and vital role [1–2]. As an integral part of China's and indeed the world's cultural heritage, classical Chinese literature and culture have always held an irreplaceable position in international Chinese language education [3–5]. These classics are not only historical testimonies but also the crystallization of the Chinese nation's spirit and cultural wisdom; they represent a unique facet of Chinese culture and embody the depth and breadth of Chinese civilization [6–8].

However, in current international Chinese language education, classical Chinese literature and culture face issues such as fragmented content and a lack of cultural depth, which hinder the comprehensive dissemination and appreciation of these works. Against the backdrop of Education 4.0, a new generation of digital technologies, represented by artificial intelligence, offers innovative solutions to these challenges. As a key cognitive intelligence technology, a knowledge graph is a structured semantic



knowledge base that describes concepts in the physical world and their interrelationships in symbolic form [9–10]. By constructing a knowledge graph for Chinese classical literature and culture, educators can better understand the full scope of this heritage, grasp the intrinsic connections between knowledge elements, and thereby conduct teaching activities in a more scientific manner [11–13]. For students, a knowledge graph of Chinese classical literature and culture can provide clear guidance on their learning path, helping them build a comprehensive knowledge system and avoid aimless study. Furthermore, the knowledge graph can intelligently recommend learning resources and suggestions tailored to students’ progress, enabling effective appreciation and personalized learning.

This paper first defines the knowledge graph construction process, adopting a top-down approach that encompasses a series of steps: identifying the resource ontology, entity recognition, knowledge extraction, and knowledge Q&A. It comprehensively employs BERT-BiLSTM-CRF to identify entity types, extracts semantic information using the BERT pre-trained model, calculates contextual information using BiLSTM and CRF layers, and predicts label sequences through the CRF layer. Entity relationships are extracted using the Bs-Spert model. A span classification module is employed to classify candidate entities, a clustering search module identifies optimal candidates, and a relationship classification module filters associations between entity pairs. After constructing the knowledge graph, a knowledge Q&A model is established to extract answer entities from the knowledge base based on the input question, thereby providing students with a tool to aid in the comprehension of classical literature.

2. Building a Knowledge Graph of Classical Chinese Literature and Culture

2.1. Knowledge Graph Construction Process

There are two primary approaches to constructing knowledge graphs: the top-down approach and the bottom-up approach. In the top-down approach, resources are first analyzed to build an ontology; data within the resources is then constrained based on the ontology; and finally, the extracted knowledge is stored in a knowledge base. In the bottom-up approach, resources are obtained from publicly available sources using technical methods, and after manual review, they are incorporated into the knowledge base. The construction process of the Chinese Classical Literature and Culture Knowledge Graph is shown in Figure 1. Given the characteristics of Chinese classical literature and culture resources, this paper adopts a top-down approach, which involves first establishing an ontology for these resources to build the model layer. Next, knowledge extraction is performed, and the extracted entities and relationships are populated into appropriate locations within the data layer. Finally, the Chinese Classical Literature and Culture Knowledge Graph is completed.

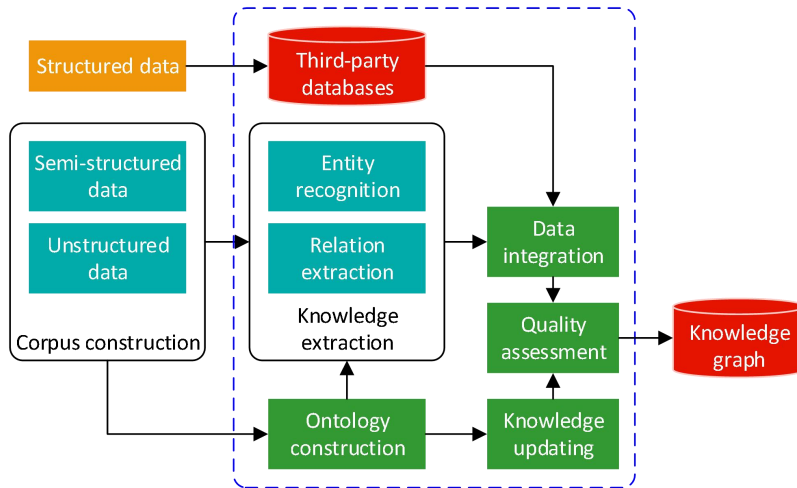


Figure 1. Knowledge graph construction process

2.2. An Entity Recognition Model Based on BERT-BiLSTM-CRF

The model described in this paper consists of BERT layers, BiLSTM layers, attention layers, and a Conditional Random Field (CRF) model. The structure of the entity recognition model is shown in Figure 2. Building upon the classic BiLSTM network model for named entity recognition, this paper incorporates a large amount of unlabeled text corpora. By utilizing the BERT pre-trained language model, we obtain word vector representations containing semantic information. We then employ attention

mechanisms to focus on key words within the text. Finally, through the CRF layer, we optimize and refine the model’s structure and training parameters, thereby completing the construction of the BERT-BiLSTM-Att-CRF model.

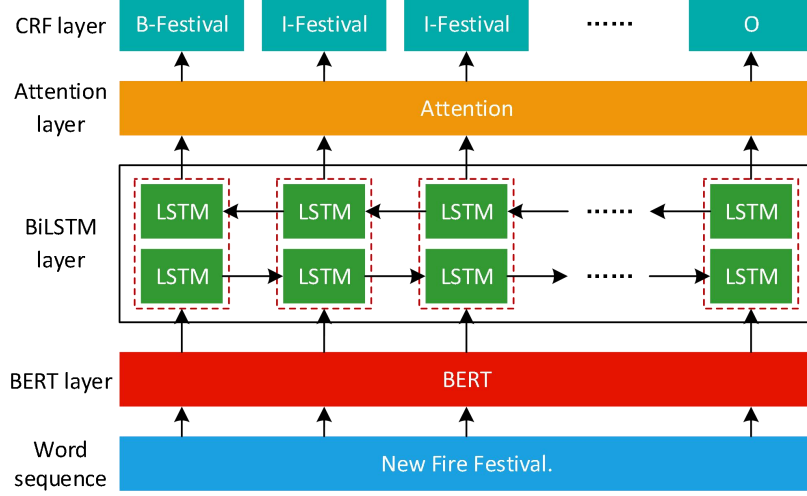


Figure 2. Entity recognition model structure

(1) BERT Layer

Since computers cannot directly process classical Chinese literary and cultural text data, the characters in the text must be converted into vector form. Traditional pre-trained language models, such as Word2vec and GPT, can both perform character vectorization. However, Word2vec achieves character vectorization by looking up a word’s representation vector in a pre-trained word embedding matrix, which is a static process. Furthermore, it cannot address the issue of polysemy; for example, the word “apple” in a sentence could refer to the fruit or to Apple Inc. The GPT model is a unidirectional language model and cannot account for the contextual information of characters or words. The BERT model can obtain dynamic word or character vectors while simultaneously resolving the issue of polysemy.

(2) BiLSTM Layer

Since the entities in the Chinese classical literature and cultural information resources have a complex structure, and entity subsequences may also be named entities—for example, the “language and text” category includes five subcategories: Chinese classical literature phonetics, vocabulary, grammar, dialects, and promotion—the relationships between words are relatively strong. Therefore, when training the model, it is necessary to fully consider the contextual information within the text sequences. However, when Recurrent Neural Networks (RNNs) process sequential data using directed recurrent structures, the issue of long-term dependencies leads to gradient explosion, preventing the model from learning features with long-term dependencies.

Although the Long Short-Term Memory (LSTM) network overcomes the gradient explosion problem found in RNNs, it can only access information preceding the current word. The Bimodal Long Short-Term Memory (BiLSTM) consists of a forward LSTM and a backward LSTM. Therefore, when performing entity recognition on Chinese classical literature cultural resources, using the BiLSTM model allows for the learning of both forward and backward information regarding characters within the text. The computational formula for BiLSTM can be expressed as:

$$f_t = \sigma(W_f h_{t-1} + U_f X_t + b_f) \quad (1)$$

$$i_t = \sigma(W_i h_{t-1} + U_i X_t + b_i) \quad (2)$$

$$o_t = \sigma(W_o h_{t-1} + U_o X_t + b_o) \quad (3)$$

$$\tilde{c}_t = \tanh(W_c h_{t-1} + U_c X_t + b_c) \quad (4)$$

$$c_t = f_t \odot c_{t-1} + i_t \odot \tilde{c}_t \quad (5)$$

$$h_t = o_t \odot \tanh(c_t) \quad (6)$$

In the formula, σ represents the activation function, f_t , i_t , and o_t represent the three gates of the LSTM at time t , and the memory cell is denoted by c_t . $U_{f,c}$ and $W_{f,c}$ represent the weight matrices corresponding to different control gates, and $b_{f,c}$ represents the bias vector. $\tilde{c}_t$ denotes the intermediate state of the input, X_t and h_t denote the input vector and output result at time t , and $\odot$ denotes the dot product operator.

The BERT-generated vector x serves as the input vector at time step t and is fed into the BiLSTM layer. By combining the feature sequences output by the forward LSTM and the backward LSTM, a concatenated vector of hidden states is obtained. After being weighted by the tanh activation function, the resulting vector h_t is fed into the Attention layer.

(3) Attention Layer

When processing contextual information, BiLSTM fails to highlight key information from the preceding text. Furthermore, in named entity recognition tasks, a single entity may have multiple representations and appear in different positions, leading to inconsistencies in annotations. Therefore, to address the characteristics of diverse naming conventions and uneven distribution of entities in Chinese classical literary and cultural information resources, the vector output by BiLSTM is fed into an Attention Layer. By utilizing the attention mechanism, the named entity recognition model reduces its focus on irrelevant information. Through the allocation of different weights, resource allocation is optimized, resulting in a more accurate feature vector. The calculation formula is as follows.

$$g_i = \sum_{j=1}^N h_j A_{i,j} \quad (7)$$

$$A_{i,j} = \frac{\exp(\text{score}(w_i, w_j))}{\sum_{k=1}^m \exp(\text{score}(w_i, w_k))} \quad (8)$$

$$\text{score}(w_i, w_j) = \frac{w_a(w_i w_j)}{|w_i| |w_j|} \quad (9)$$

Here, $A_{i,j}$ represents the weight, h_j represents the output of the previous-layer BiLSTM, w_i denotes the current word, and w_j denotes the words in the document containing the current word. $\text{score}(w_i, w_j)$ is used to calculate the similarity between two words by computing the cosine distance between them. w_a represents the parameters, which are learned during training.

(4) CRF Layer

The BiLSTM and attention layers only take contextual information into account; they do not consider the dependencies present within the labels. Furthermore, the BiLSTM layer outputs the label with the highest score for a given unit and uses it as that unit's label. Although the resulting labels are all correct, the BiLSTM and attention layers cannot guarantee that every label prediction will be correct; therefore, the final character label cannot be determined based solely on the results from these layers. The Conditional Random Field (CRF) model uses a transition matrix to capture dependencies between labels and obtain the globally optimal sequence of labels. The probability formula for the input label sequence x and the output label sequence y is:

$$s(x, y) = \sum_n \left(\sum_{i=1}^m P_{i, y_i} + \sum_{i=0}^m A_{y_i y_{i+1}} \right) \quad (10)$$

Here, $A_{y_i y_{i+1}}$ represents the transition score between labels, and p_{i, y_i} is the probability that the current character is predicted to be a label. Then, $\text{Soft max}(\cdot)$ is used to calculate the probability of the label sequence y . Finally, the Viterbi algorithm is used to compute the final score for y , and the label sequence with the highest score is selected as the model's output.

2.3. Entity-Relation Extraction Based on Bs-Spert

This paper employs the BERT pre-trained model as the foundation for the Bs-Spert model to conduct research on the joint extraction of entities and entity relationships in classical Chinese literature and culture. The Bs-Spert model described in this paper classifies spans across all entities to identify them; its main components consist of the BERT pre-trained model module, the clustered search module, the span filtering module, and the relationship classification module.

2.3.1. Span Classification Module

Any possible entity span candidates can be fed into the span classifier. Here, we first assume that the input span sequence is $s = \{q_i, q_{i+1}, \dots, q_{i+k-1}, q_{i+k}\}$ and the span length is $k+1$. The predefined entity set is $\mathcal{V} \cup \{O\}$, where $\mathcal{V}$ represents all existing, identifiable text entity categories, and O represents entity spans that do not exist or cannot be identified.

Span embeddings: The span embeddings obtained from the pre-trained BERT model are combined into $f(q_i, q_{i+1}, \dots, q_{i+k-1}, q_{i+k})$ using the fusion function f . This paper thoroughly investigates the impact of different fusion functions on the final experimental results and ultimately finds that the highest experimental accuracy is achieved using the max-pooling function.

Span-length embedding: For the span sequence s given above, the span length is $k+1$. We then look up the embedding w_{k+1} with span length $k+1$ from a dedicated matrix containing each fixed span length (span lengths of 1, 2, 3, 4, etc.), where the span-length embeddings are learned via backpropagation and allow the model to incorporate prior knowledge about span lengths. Concatenating the span length embedding with $f(q_i, q_{i+1}, \dots, q_{i+k-1}, q_{i+k})$ yields a span representation of width $k+1$:

$$q(s) = f(q_i, q_{i+1}, \dots, q_{i+k-1}, q_{i+k}) \quad (11)$$

Sentence vector c : Sentence vector c is obtained from the pre-trained model BERT. This vector represents the historical contextual information of the sentence text and functions similarly to keywords within a sentence. Additionally, by leveraging this contextual history, sentence vector c effectively resolves entity ambiguity issues. Therefore, the use of sentence vector c is indisputable, and it plays a crucial role in improving accuracy within the input composition of the span classification module. Finally, the input to the span classification module after concatenating the sentence vector c is shown in Equation (12):

$$x^s = q(s) \cdot c \quad (12)$$

Finally, x^s is fed into the span classification module, which primarily uses the softmax function as the classifier. As shown in Equation (13), this generates posterior probabilities for each entity class, including, of course, the posterior for O (spans for entities that do not exist or cannot be identified):

$$y_s = \text{softmax}(W^s \cdot x^s + b^s) \quad (13)$$

Here, W^s represents the weight matrix for the input x^s , and b^s represents the bias term.

2.3.2. Clustered Search Module

The Beam Search algorithm, also known as beam search, is a heuristic search algorithm that can be viewed as a compromise between greedy algorithms and exhaustive search. Beam search can be viewed as an improvement over greedy search. Compared to greedy algorithms, beam search expands the search space. However, in cases where the solution space is relatively large, beam search prunes some lower-quality nodes during each step of depth expansion, retaining only higher-quality nodes. Nevertheless, the search space of beam search is far smaller than the exponential search space of exhaustive search. In Beam Search, there is only one parameter: Beam Width. Here, it is set to k , which selects the k optimal results with the highest conditional probabilities under the current conditions as the first words of the candidate output sequences. Based on the current word, the k optimal results with the highest conditional probabilities among all selected combinations are chosen. Throughout the selection process, a pool of k optimal candidates is maintained, and the final optimal result is selected from all candidates.

2.3.3. Relationship Classification Module

In this section, the set of classical literary entities is defined as R based on the annotated

relationships between entities. The relationship classification module is used to determine whether candidate entity pairs obtained from the span filtering module belong to the classical literary entity set R . The input to the relationship classification module primarily consists of two components: a $q(s)$ formed from the span embeddings and span lengths processed by the fusion BERT model, and the contextual history information between the two valid entities.

Another component of the input to the relationship classification module is the contextual history information between entity pairs (s_1, s_2) . Since this involves obtaining contextual history information, it is natural to mention the sentence vector c obtained through the BERT model. However, in the early stages of the experiment, using the simple sentence vector c as the input for contextual history information resulted in unstable experimental results; specifically, when dealing with scenarios involving single relationships in classical Chinese literary texts, the sentence vector c performed quite well. As is well known, natural language texts do not merely contain simple scenarios involving single relationships; more often, a single sentence contains multiple complex relationships. When dealing with such complex texts, simply using the sentence vector c as the contextual input inevitably proves inadequate. Therefore, in this paper, we use the BERT pre-trained model to obtain range embeddings $d(s_1, s_2)$ between valid entity pairs via pooling functions. We then feed $d(s_1, s_2)$ as contextual information into the relationship classification module. However, there are exceptions to every rule: if entity overlap occurs, it can lead to the anomalous situation where the range between valid entity pairs is empty. Consequently, we define the range embedding between valid entity pairs as $d(s_1, s_2) = 0$.

Next, we perform a concatenation operation similar to that used in the span classification module, combining the span embeddings of the two valid entity pairs with the range embeddings between the entity pairs. At the same time, we account for the asymmetry in the relationships between the entity pairs. Consequently, the input representation in the relationship classification module is shown in Equation (14):

$$\begin{cases} x_1^r = q(s_1) \cdot d(s_1, s_2) \cdot q(s_2) \\ x_2^r = q(s_2) \cdot d(s_1, s_2) \cdot q(s_1) \end{cases} \quad (14)$$

Input x_1^r and x_2^r into the single-layer classifier in the relationship classification module, as shown in Equation (15) below:

$$\begin{cases} y_1^r = \sigma(W_1^r \cdot x_1^r + b_1^r) \\ y_2^r = \sigma(W_2^r \cdot x_2^r + b_2^r) \end{cases} \quad (15)$$

Here, W_1^r and W_2^r represent the weight matrices for x_1^r and x_2^r , respectively, while b_1^r and b_2^r represent the corresponding bias terms. σ denotes a sigmoid function with a dimension size of R ; any high response observed in the sigmoid layer indicates that the relevant entities have a corresponding relationship with (s_1, s_2) . In this paper, we set a response threshold of α . It stated that when the response score for any relationship is greater than or equal to α , it can be assumed that a known relationship $r \in R$ exists between the two valid entity pairs; conversely, if the response score is less than the response threshold α , it can be assumed that no known relationship $r \notin R$ exists between the two valid entity pairs.

3. A Knowledge-Based Question-Answering Model Using Knowledge Graph Embeddings

A knowledge graph- or knowledge base-based question-answering system is defined as one that extracts answer entities from a knowledge base in response to a given input question. Knowledge is typically stored in a knowledge base in the form of triples, with each fact representing a relationship between two entities. Each triple connects two entities via a relationship. Simple questions constitute the majority of queries in knowledge graph embedding-based question-answering; if the correct head entity and relationship are identified, any tail entity can serve as the answer. Entities and relationships in the knowledge graph are first embedded into two low-dimensional spaces using knowledge representation learning methods, such that each fact (h, r, t) can be represented by three latent vectors (e_h, r_i, e_t) . Therefore, given a question, the model only needs to predict the corresponding facts e_h and r_i —that is, locate them in the embedding space of entity relationships—to obtain the correct answer.

3.1. Knowledge Graph Embedding

Knowledge graph embedding creates a vector representation for each node and relationship in a knowledge graph within a low-dimensional, dense vector space by preserving the original connections and relationships in the graph. The core idea can be summarized as follows: a knowledge graph embedding model predicts the tail entity e_t given a head entity and a relationship using a scoring function $f(e_h, r)$. Initially, the vector embeddings for entities and relationships are initialized randomly or using trained word vectors, and the model learns by minimizing the discrepancy between the actual embeddings and the predicted values. Traditional knowledge graph embedding models generally perform well when predicting one-to-one relationships, but when faced with one-to-many, many-to-one, or even many-to-many relationships, they perform poorly because the granularity of the knowledge graph embeddings is too coarse to adequately model entity and relationship representations. To address this issue, this paper employs a knowledge embedding method to perform embedding representations on the Chinese Classical Literature and Culture Knowledge Graph.

3.2. Header Entities, Relation Learning, and Answer Computation

Given a simple question, the goal of this paper is to find a point or vector in the relation embedding space that serves as the feature representation r for the relation in the question, and to find a point or vector in the entity embedding space that serves as the feature representation e_h for the head entity. For any question that a knowledge graph can answer, its relationship vector r must exist in the relationship embedding space. To this end, the goal of this section is to take a question as input and return a vector that represents the relationship embedding of that question as closely as possible.

First, the sentences are tokenized, and their word vectors are fed into two separate forward and backward LSTMs. The hidden state of each token is represented by concatenating the corresponding forward and backward hidden states, denoted as $h_i = (\vec{h}_i; \overleftarrow{h}_i)$. The attention weights for each hidden state h_i are calculated as shown in Equations (16) and (17):

$$\alpha_i = \frac{e^{q_i}}{\sum_{j=1}^L e^{q_j}} \quad (16)$$

$$q_i = \tanh(w^T [x_i; h_i] + b_q) \quad (17)$$

Using attention weights α_i , a new representation $s_i = [x_i; \alpha_i h_i]$ is created for each token and fed into the fully connected layer to predict the representation $(r_i \in \mathbb{R}^{d \times 1})$ for each token. The predicted representation is obtained by averaging r_i vectors. As shown in Equation (18):

$$\hat{p}_n = \frac{1}{N} \sum_{i=1}^N r_i^T \quad (18)$$

Next, we proceed to the head entity learning stage. The word embeddings from the word segmentation step are fed into a bidirectional LSTM model to generate the corresponding hidden state h_i . The hidden state is then passed to a fully connected layer, which produces a vector $v_i \in \mathbb{R}^{2 \times 1}$ via a Softmax layer. This layer extracts the identified potential head entities, forming the set of candidate entities.

Finally, if the head entity of a triplet belongs to the set of candidate entities, this paper refers to it as a candidate triplet. By measuring the distance between the candidate triplet and the predicted representation, we identify the closest predicted entity as the answer and set the distance as the confidence level of the correct answer.

4. Pilot Test of the Knowledge Graph for Classical Chinese Literature and Culture

4.1. Knowledge Graph Entity Naming Test

The data used for model training in this paper is primarily sourced from Minzu.cn, Baidu Baike, Wikipedia, and specialized books covering the field of classical Chinese literature and culture. This paper uses precision (P), recall (R), and F1 score to evaluate the performance of entity recognition. The model performs entity recognition on 12 annotated entity categories within the corpus of classical Chinese

literature and culture. The entity recognition results of the BERT-BiLSTM-CRF model for the knowledge graph are shown in Table 1. Among the 12 entity types in the Chinese classical literature and culture knowledge graph, Chinese characters achieved the highest recognition accuracy at 91.89%. Works achieved the highest recall at 90.94%, while events achieved the highest F1 score at 92.89%. The BERT-BiLSTM-CRF model demonstrates good adaptability in the field of Chinese classical literature and culture.

Table 1. The entity recognition results of the knowledge graph

Entity type	P (%)	R (%)	F1 (%)
Chinese characters	91.89	83.95	88.69
Vocabulary	90.26	83.44	88.19
Grammar	84.28	82.20	87.60
Cultural Knowledge Points	85.41	87.67	86.14
Characters	80.53	77.62	81.47
Location	89.14	86.69	86.08
Event	80.01	78.59	92.89
Composition	83.63	90.94	88.30
Genre	79.58	89.06	79.81
Image	75.27	89.68	80.26
Bibliographic metadata	75.99	76.54	86.53
Object	86.41	75.80	83.90

This paper compares the LSTM+CRF model, the BERT+CRF model, the Lattice LSTM model, the SoftLexicon model, the ATT-CRF model, the Bi-Transformer model, the GPT model, and the BERT-BiLSTM-CRF model proposed in this paper. Among these, the first two models use purely character-level inputs without incorporating lexical information, while the remaining models do incorporate lexical information. The knowledge graph entity recognition results for the different models are shown in Table 2. The performance of the BERT+CRF model differs little from that of the LSTM+CRF model, with a precision difference of only 1.95%. This is primarily because the CRF’s powerful contextual encoding capability enables it to extract the necessary information, whereas BERT merely selects relevant information from the LSTM for processing, resulting in a negligible improvement. In contrast, the BERT-BiLSTM-CRF model proposed in this paper achieved accuracy, recall, and F1 scores of 92.65%, 93.63%, and 90.47%, respectively, outperforming all other models in all metrics.

Table 2. The entity recognition results of knowledge graphs from different models

Method	P (%)	R (%)	F1 (%)
LSTM+CRF	84.82	78.63	82.45
BERT+CRF	86.77	77.59	83.62
Lattice LSTM	89.12	88.15	89.03
SoftLexicon	89.94	90.34	89.24
ATT-CRF	89.22	89.24	90.81
Bi-Transformer	89.18	89.50	88.67
GPT	90.95	88.34	88.30
BERT-BiLSTM-CRF	92.65	93.63	90.47

4.2. Knowledge Graph Relationship Extraction Test

In this section, we report the main experimental results and analyze the performance of the proposed Bs-Spert model. Furthermore, by comparing the experimental results with those of baseline methods, we provide a more in-depth analysis of our work and the experimental findings. The PR curves for relation extraction from the Chinese Classical Literature Knowledge Graph using different models are shown in Figure 3. The Bs-Spert model designed in this paper demonstrates superior performance compared to other extraction methods. Once the recall exceeds 0.15, the PR curves of the Bs-Spert model show a stable trend; in particular, when the recall exceeds 0.25, the Bs-Spert model’s precision significantly outperforms that of the other models. This indicates that the proposed model has certain advantages in predicting different types of relationship categories.

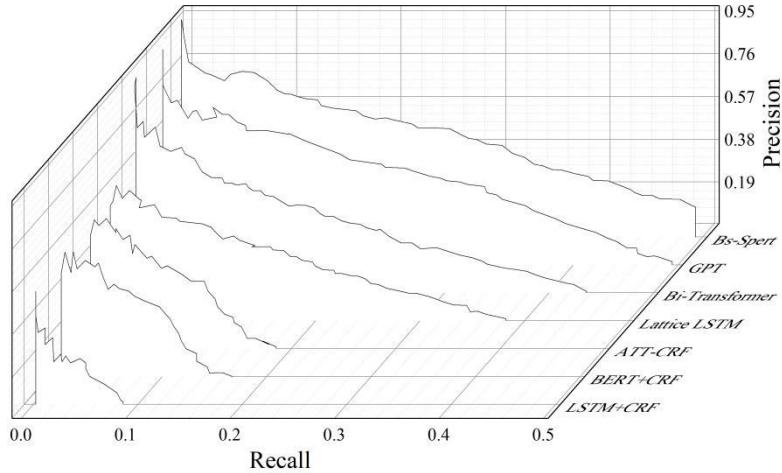


Figure 3. The PR curve of relation extraction in knowledge graphs

This paper uses all instances in the dataset for model evaluation; the $P@N$ experimental results are shown in Table 3. The results indicate that the Bs-Spert model designed in this paper significantly outperforms all baseline methods on most metrics. Overall, the Bs-Spert model presented in this paper demonstrates excellent accuracy at $P@150$, $P@300$, $P@450$, $P@600$, $P@750$, and $P@900$, as well as in the average accuracy across these six values. It achieves high accuracy at $P@150$, $P@300$, $P@450$, $P@600$, $P@750$, and $P@900$, it outperformed the suboptimal GPT model by 9.02%, 1.93%, 6.26%, 8.86%, 4.66%, and 2.44%, respectively. The results demonstrate that the model proposed in this paper performs remarkably well across the entire distribution. Furthermore, compared to all baseline models, the model in this paper achieves the highest average $P@MEAN$ across the six evaluation metrics. The results indicate that the model exhibits strong robustness when handling noisy instances by combining local and global structural information within sentences.

Table 3. The $P@N$ experimental results of relation extraction

Method	$P@150$	$P@300$	$P@450$	$P@600$	$P@750$	$P@900$	Mean
LSTM+CRF	59.87	54.41	48.10	47.24	41.10	45.70	49.40
BERT+CRF	60.91	56.05	49.79	50.52	51.45	46.05	52.46
ATT-CRF	67.54	58.56	56.97	51.18	55.59	51.46	56.88
Lattice LSTM	71.37	64.34	58.49	55.54	57.19	57.11	60.67
Bi-Transformer	73.33	66.77	59.05	59.22	60.75	57.85	62.83
GPT	77.46	78.15	59.30	60.50	65.52	60.43	66.89
Bs-Spert	85.14	79.69	63.26	66.38	68.74	61.94	70.86

4.3. Question-Answering Test in the Knowledge Graph Domain

Using the metrics EM, WER, and MER, we compared the knowledge Q&A model embedded in the knowledge graph described in this paper with other models. The performance comparison of the knowledge Q&A models is shown in Table 4. Here, EM represents the match rate of knowledge Q&A, while WER and MER can be used to evaluate the overall performance of candidate-based information extraction models on mixed data. WER effectively assesses a model’s ability to identify incorrect information, and a lower WER value indicates that the model has extracted more valid information; MER effectively evaluates the model’s ability to obtain complete valid information; similarly, a lower MER value indicates that the model has extracted more comprehensive valid information.

The knowledge Q&A model in this paper achieved the best results on the Chinese Classical Literature and Culture Knowledge Graph, with a computational match rate of 58.81%, demonstrating that the candidate extractor can effectively extract valid information relevant to the question from mixed datasets. Specifically, the two multi-head attention modules in the candidate extractor can respectively obtain multi-level semantic information relevant to the question from text and tables; the source selector in the inference layer can correctly locate the positions of candidate fragments within the mixed dataset.

Table 4. Performance Comparison of Knowledge Q&A

Method	EM (%)	WER (%)	MER (%)
BERT-RC	10.09	53.47	54.83

NumNet+V2	37.25	50.67	53.84
TaPas for WTQ	19.31	48.49	53.47
HyBridr	6.62	47.14	44.31
TAGOP	55.51	36.91	33.85
KGQA+RAG	54.74	35.69	33.67
KARPA	55.78	34.89	33.44
Ours	58.81	33.07	33.32

5. Conclusion

This paper adopts a top-down approach to construct a knowledge graph of Chinese classical literature and culture. It employs a BERT-BiLSTM-CRF entity recognition model and uses Bs-Spert to extract relationships between entities. By implementing knowledge graph embedding for knowledge-based question-answering, the model assists students in international Chinese language education with the appreciation of classical literature. The research conclusions are as follows:

(1) Among the 12 entity types in the Chinese classical literature knowledge graph, the BERT-BiLSTM-CRF model achieved a recognition accuracy of 91.89% for “Chinese characters” and a recall rate of up to 90.94% for “works.” Overall, the model achieved an accuracy of 92.65%, a recall rate of 93.63%, and an F1 score of 90.47%.

(2) The PR curve of the Bs-Spert model for cultural knowledge extraction from Chinese classical literature demonstrates a significant advantage over the curves of other models. It is capable of fully extracting both local and global structural information and exhibits good robustness across different noisy instances.

(3) Knowledge graph-based question-answering in the field of classical literature and culture also demonstrates superior performance metrics, with an answer match rate as high as 58.81%, significantly higher than that of other models.

The knowledge graph developed in this paper for Chinese literature performs well in all aspects. By leveraging this knowledge graph, we can uncover the essence of China’s outstanding traditional culture and enhance the understanding of traditional culture among learners of Chinese as a foreign language.

References

1. Fengrong, J. (2023). Promoting the International Dissemination of Chinese Culture Through International Chinese Language Education: A Case Study of Chinese-English Idiomatic Equivalence. *Contemporary Social Sciences*, 8(6), 119.
2. Ying, W. (2016). Review of the Confucius Institutes’ strategy for the dissemination of Chinese culture. *Chinese Education & Society*, 49(6), 391-401.
3. Zhang, B., & Haoxuan, Z. (2025). The study of classic Chinese literature from the perspective of historical poetics. *Neohelicon*, 52(1), 77-89.
4. Peng, J. (2014). On aesthetic conception of the ontological foundation of classical Chinese literature. *Asian Journal of Social Sciences & Humanities Vol*, 3, 4.
5. Su, Y. (2022). The Influence of Ancient Chinese Cultural Classics in Southeast Asia. In *A Study on the Influence of Ancient Chinese Cultural Classics Abroad in the Twentieth Century* (pp. 37-58). Singapore: Springer Singapore.
6. Tao, Z. (2023). The Embodiment of “Harmony” Culture in the Classics of Modern and Contemporary Chinese Literature. *Journal of Humanities, Arts and Social Science*, 7(9).
7. Zhang, Q. (2024). Studying the Ancient Classics to Meet Present Needs: The Spirit of Shouldering Responsibility for the Whole World. In *The Essence of Chinese Humanistic Culture* (pp. 89-109). Singapore: Springer Nature Singapore.
8. Zhuang, P. (2021). On Variations of Classical Chinese Literary Theory for a Framework of Global Literary History. *Cultura*, 18(1), 23-40.
9. Hao, X., Ji, Z., Li, X., Yin, L., Liu, L., Sun, M., ... & Yang, R. (2021). Construction and application of a knowledge graph. *Remote Sensing*, 13(13), 2511.

10. Zou, X. (2020, March). A survey on application of knowledge graph. In *Journal of Physics: Conference Series* (Vol. 1487, No. 1, p. 012016). IOP Publishing.
11. Cui, Y., Yao, S., Wu, J., & Lv, M. (2024). Linking past insights with contemporary understanding: an ontological and knowledge graph approach to the transmission of ancient Chinese classics. *Heritage Science*, 12(1), 1-18.
12. Wei, Y., Wang, H., Zhao, J., Liu, Y., Zhang, Y., & Wu, B. (2020, August). GeLaiGeLai: a visual platform for analysis of classical Chinese poetry based on knowledge graph. In *2020 IEEE International Conference on Knowledge Graph (ICKG)* (pp. 513-520). IEEE.
13. Guo, W. (2022). Correlation between the dissemination of classic english literary works and cultural cognition in the new media era. *Advances in Multimedia*, 2022(1), 3616432.