

Research on the Construction and Application of Educational Big Data-Driven Subject Knowledge Graph in Colleges and Universities

Na Guo¹, 

¹ School of Information Engineering, Institute of Disaster Prevention, Sanhe, Hebei, 065201, China

ABSTRACT

In the long-term teaching practice, various disciplines have accumulated a large number of teaching resources but cannot function fully and efficiently. For this reason, this study constructs a knowledge mapping of college disciplines based on deep learning. First of all, the overall construction of the atlas is planned, the core concepts of the discipline are identified, the relationships between the knowledge points are defined, and the resources corresponding to the knowledge entities and attributes are expanded. Then deep learning is utilized for the entity construction of the subject knowledge graph, the neural network models BiLSTM+CRF and BiLSTM+Attention are used for the subject entity identification and relationship extraction, and finally the subject knowledge fusion and storage is carried out, and the effectiveness of the designed algorithms is verified on the dataset. The data show that the knowledge representation of knowledge graph is conducive to demonstrating the logical meaning between learning materials, facilitating learners to correlate what they have learned previously with what they are learning now, fusing old and new knowledge, and facilitating learners to meaningfully construct knowledge.

Keywords: Knowledge graph, BiLSTM, conditional random field, knowledge fusion

1. Introduction

Knowledge graph is the new form and new product of knowledge engineering in the current period of rapid development of artificial intelligence and big data technology, and is now developing with other emerging information cross-fertilization [1]. The Language and Knowledge Computing Specialized Committee of the Chinese Language and Information Society of China has reviewed the development history of knowledge engineering over the past forty years, and found that knowledge engineering consists of five iconic phases, the pre-knowledge engineering period, the expert system period, the World Wide Web 1.0 period, the group intelligence period, and the knowledge graph period [2-3].

 Corresponding author.

E-mail address: gn18810866177@163.com (N. Guo).

Received 10 March 2024; Revised 15 May 2024; Accepted 05 December 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-224](https://doi.org/10.61091/jcmcc127b-224)

© 2025 The Author(s). Published by Combinatorial Press. This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by/4.0/>).

Among them, there are several key events and concepts in the development history of knowledge graph, which can be used as the key features to divide the technological development history of knowledge graph [4].

Knowledge graph, as the core driving force to promote the development of artificial intelligence with efficient semantic processing, is a tool to manage and analyze modern knowledge and provide decision support [5-6]. With entity concepts as nodes and relationships as edges, it can reveal knowledge and the complex associative relationships between them in an intuitive and visual way, realizing a clear graphical representation of the structural relationships of knowledge [7-8]. Knowledge mapping can expand new ideas for the innovation of teaching mode in the new situation and provide the necessary tools for knowledge system construction [9]. In the field of education, knowledge mapping can start from a large number of disordered educational information resources, reconfigure the connection between knowledge, effectively organize the knowledge system of various disciplines, and then assist students to efficiently and accurately learn the weak knowledge points, make up for the missing fragments of their knowledge system, and rapidly improve their disciplinary knowledge system [10-12]. In addition, with the help of knowledge mapping, through the analysis of students' individual known ungraspable knowledge points, to find the implicit unintelligible related knowledge points, to provide students with targeted exercises recommended to mobilize the subjective initiative, reduce the burden of homework at the same time to achieve precision, personalized customized teaching. Knowledge mapping can also help teachers to establish subject knowledge lecture system more quickly and orderly and assist educational experts to evaluate subject examination papers, etc. [13-15].

Driven by the wave of informationization and artificial intelligence, the field of higher education presents the trend of diversification of teaching technology and decentralization of educational resources. Along with the decentralization of teaching resources, problems such as knowledge fragmentation and the lack of effective interdisciplinary integration between disciplines have gradually emerged [16-17]. Discipline knowledge mapping technology is based on knowledge mapping technology, which can effectively deal with massive, complex and dispersed discipline knowledge and realize the management and utilization of discipline knowledge. Currently, most of the methods of disciplinary knowledge mapping construction are proposed based on the knowledge mapping construction of a single course or a single discipline. And nowadays, the cross integration of disciplines is a general trend, and multidisciplinary intersection puts forward new requirements for the integration of interdisciplinary knowledge in the construction of disciplinary knowledge maps [18-20].

The innovative application of big data, artificial intelligence and other technologies in the field of education has prompted the development of education informatization in the direction of personalization, intelligence and refinement. The return of the new generation of artificial intelligence technology with deep learning and knowledge mapping as the core will reshape and reengineer the personalized adaptive learning system [21-23].

Knowledge mapping is an important application of visualization technology in bibliometrics, which has been widely distributed with us. The key core points of knowledge mapping construction technology mainly contain automated and intelligent mining of textual knowledge, semantic connectivity establishment and unified standards for data structuring. Chen, P et al. designed a KnowEdu system with the function of automatic and automated construction of educational knowledge mapping, and combined with the practical cases to verify the feasibility of the system, confirming that the designed KnowEdu system is feasible and effective [24]. Zhang, N et al. elaborated that the emergence of big data technology creates favorable conditions for the construction of knowledge chains in knowledge graphs, so based on the four principles of Linked Data Publishing Content Objects and their semantic features, and using the RDF data model, the unstructured data as well as structured data with different standards are converted in a unified manner and their

associations are established to form an intelligent, semantic knowledge graph [25]. Melnyk, I et al. introduced an end-to-end knowledge graph generation framework, which mainly includes graph generation nodes based on pre-trained language models and efficient text mining based on simple edge construction headers, and the knowledge graph generation framework was put into the WebNLG 2020 Challenge dataset for evaluation and testing, which confirmed that the proposed knowledge graph generation system has good performance [26].

The main application areas of knowledge graph contain subject knowledge organization and aggregation, teaching support, and students' personalized learning, etc. Deng, Y established and evaluated the Massive Open Online Course (MOOC)-based Knowledge Graph for Higher Education, and pointed out that the Knowledge Graph can effectively support learners' personalized learning, and the personalized learning plan developed based on the Knowledge Graph provides learners with targeted and enriched teaching services [27]. Yang, H et al. envisioned a multimodal domain knowledge mapping strategy for the integration of decentralized knowledge in computer science, which effectively enables the mining and establishment of semantic relationships between multimodal data [28]. Nitchot, A et al. envisioned a knowledge mapping structure for aiding teaching and learning as well as other pedagogical aids and targeted performance evaluations that concluded that these educational knowledge aids facilitate knowledge assimilation and higher learning achievement [29]. Zou, X et al. attempted to construct a knowledge map related to road safety and therefore systematically reviewed road safety related literature studies, focusing on analyzing and collecting the knowledge base, topic distribution, research frontiers, and research trends of road safety research, and visualizing the mining and integration [30]. The research and development of knowledge graph has been more mature, and its future optimization path is mainly for the research of knowledge graph refinement. Paulheim, H elaborated the current state of practice of free knowledge graph and commercial knowledge graph, and argued that these semi-structured knowledge graphs oscillate between completeness and correctness, thus attempted to explore the research literature on the path of knowledge graph refinement, and the research literature on these proposed optimization strategies are critically considered and evaluated [31].

In this paper, we first propose the construction process of disciplinary knowledge graph, and set up the knowledge entity, object attributes and disciplinary resource management. Then according to the three-phase ten-process method for the construction of discipline ontology, the characteristics of the discipline are analyzed, knowledge points are collected and processed by crawling the related content, the dataset for entity recognition is constructed, and the BiLSTM+CRF algorithm for entity recognition and the BiLSTM+Attention algorithm for entity extraction are proposed. Finally, the data from MySQL and Ownthink knowledge base are merged through knowledge fusion to realize knowledge graph visualization and intelligent Q&A. It is applied with a university teaching to observe the effect.

2. The process of building a disciplinary knowledge map

2.1. *Disciplinary Knowledge Mapping Steps*

The construction level of disciplinary knowledge mapping is divided into three main issues: knowledge entity setting, object attribute setting and disciplinary resource management.

1) Setting of knowledge entities: there are many kinds of disciplinary knowledge entities, which should be streamlined in the construction of the type and number of knowledge entities, covering the main core categories. The setting of knowledge entities requires two types of specific work: summarization and organization of knowledge points and selection of core concepts. First, based on the syllabus and teachers' experience, the subject knowledge points with teaching significance are summarized and organized, and these knowledge points are the preliminary knowledge entities of the knowledge map. Second, the knowledge entities are categorized, and the core concepts of the discipline are categorized and identified. In the categorization process, if a certain type of knowledge entity is too little and at

the same time more important, the knowledge entity will be merged into a more closely related category. If there are few and unimportant knowledge entities in a category, they are eliminated. In this process, we should follow the principle of generality in the construction of knowledge maps of subject categories, and also adjust dynamically according to expert opinions.

2) Setting of Object Attributes: The association between subject knowledge entities is particularly rich, with link structure (hierarchical structure) and mesh structure (associative relationships), and appropriate object attributes should be selected for combing and describing these relationships during the construction of the atlas. The setting of object attributes has two key aspects: the determination of object attributes and the setting of object attributes between different entities.

The setting of object attributes between different entities relies on expert experience. In this session, based on the determined knowledge entities and object attributes, the knowledge entities are compared two by two to set appropriate object attributes. Multiple object attributes can exist between two knowledge entities or no object attributes can exist. The work is mainly done manually and reviewed by experts.

3) Discipline resource management: describing knowledge points in the form of long text and managing discipline resources in the form of resource links. At present, the discipline monographs use long text to describe the knowledge points, and there is no better new method here, so it continues to be used. However, it is possible to correspond resource links to subject knowledge points for management and effective utilization. The specific realization method is to upload the pictures, videos or tables corresponding to a certain knowledge entity, upload them to the cloud server, obtain the discipline resource links, and set the resource links as the data attributes of the knowledge entity.

To summarize, disciplinary knowledge mapping should contain diverse knowledge entities of disciplines, clearly reflect the hierarchical structure of knowledge points and associated relationships, and at the same time describe the knowledge points in long text and manage the disciplinary resources in the form of links.

2.2. Academic Knowledge Graph Construction Process

Prior to the formal construction work, the disciplinary knowledge mapping construction process is designed to provide guidance for subsequent mapping implementations.

2.2.1. Defining the initial process. Disciplinary knowledge mapping is often constructed with reference to a five-step process versus a seven-step process. The five-step process includes: (1) delineating the domain ontology and scope; (2) acquiring disciplinary knowledge; (3) enumerating important knowledge concepts; (4) determining the hierarchy and attributes of knowledge concepts; and (5) coding materialization and resource filling.

However, neither the five-step method, the seven-step method, nor the reference process can be effective in solving the three important problems of disciplinary knowledge mapping construction:

1) All the above three methods are limited to the study of knowledge entities and relationships in a certain domain or a certain discipline, focusing on the realization of the knowledge graph, without considering the applicability of the knowledge graph.

2) The five-step and seven-step processes are more effective in constructing knowledge graphs with a relatively single knowledge entity, and the construction can be completed according to the process. The disciplinary knowledge mapping is an applied discipline with multi-disciplinary knowledge intersection, which needs to determine the core concepts by combining the disciplinary characteristics. And the five-step and seven-step methods have no corresponding process.

3) Discipline is essentially an overall description of the knowledge points of the discipline in the teaching scene, so the relationship between the knowledge points of the discipline is complex, and all kinds of relationships are mixed together, which makes it difficult to define the relationship between the knowledge points of the discipline. If only in accordance with the five-step and seven-step process,

when defining the attributes between concepts, in the face of the complex relationships between disciplines, there is still no way to start, and the original process lacks guiding significance, so it is necessary to combine with the characteristics of the discipline to split and refine this key process.

In response to Problem 1, the construction of disciplinary knowledge mapping is divided into three stages: the preparation stage, the system construction stage and the realization and application stage. In the preparation stage, the overall planning of the construction of the map is carried out by combining the application scenarios of the disciplinary knowledge map in colleges and universities. In the realization and application stage, suitable storage tools will be selected around the application and secondary development of the knowledge graph.

For problem 2, after summarizing the subject knowledge and before defining the concept hierarchy, a process is added to identify the core concepts of multiple categories of the subject.

In response to question three, the difficult task of defining inter-conceptual attributes is split up. Firstly, the objective categorization relationship between knowledge points is determined according to the hierarchical structure of the subject; secondly, the production-related relationship between knowledge points is determined according to the links in the process; and finally, the learning sequence relationship between knowledge points is determined according to the actual pedagogical needs of the subject. Therefore, the original "Define inter-conceptual attributes" was split into "Define objective categorical relationships contained in the subject" and "Define production-type relationships contained in the subject", "Defining the pedagogical relations contained in the subject".

According to the above formulation, the reference process was modified to obtain a three-stage ten-process approach for the construction of disciplinary knowledge maps as shown in Fig. 1:

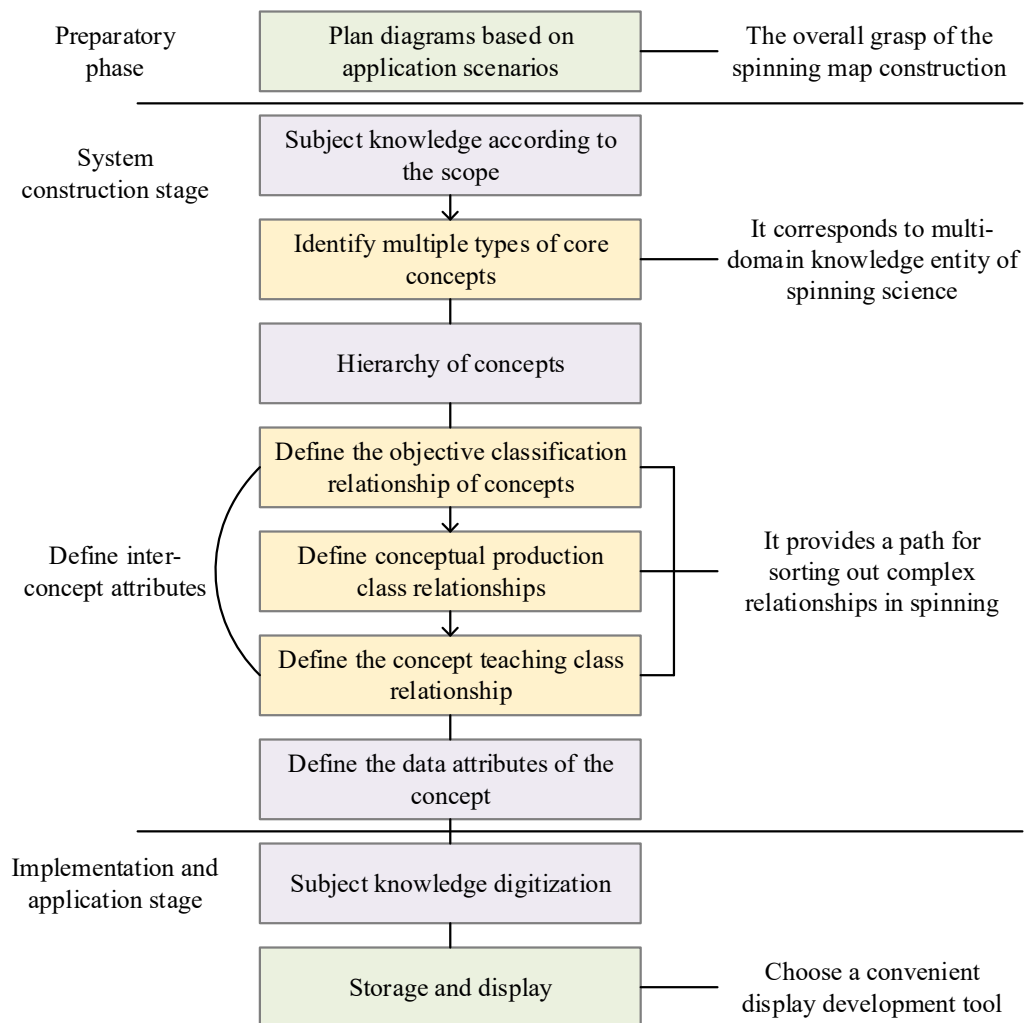


Fig. 1. Three-stage ten-process method for constructing knowledge map

The method is a combination of processes proposed after combining the five-step and seven-step processes as well as the characteristics of the discipline, which is more adaptable to the construction of disciplinary knowledge graphs compared to the five-step and seven-step methods.

2.2.2. Selection of technologies for each process. The previous section identifies the process of disciplinary knowledge graph construction, and at the level of specific process operation, it is necessary to choose appropriate methods or tools in combination with the characteristics of the discipline. The methods and tools involved in the construction of knowledge graphs will be unified here, and the following is an analysis of the scope of application of common disciplinary knowledge graph construction tools from different aspects such as construction principles, construction methods, construction processes, ontology construction tools, databases and disciplinary knowledge acquisition methods.

1) Construction principles: large data size, high degree of data resource structuring, single mapping relationship mapping

Suitable for bottom-up construction principle: maps with small data scale, low degree of data resource structuring, and complex mapping relationships are suitable for top-down construction principle.

2) Construction methods: manual construction methods are suitable for disciplines with complex structure, complex resources and small data scale; crowdsourcing methods are suitable for disciplines with complex structure and large data scale; automated construction methods are suitable for disciplines with simple structure and large data scale.

3) Ontology construction tools: Protégé is suitable for Chinese data that is unstructured and with small data size, and Jena is suitable for data that is structured and with large data size [32].

4) Databases: databases are divided into relational databases and graph databases, graph databases are suitable for large-scale data association query and multi-hop query, and relational databases are suitable for structured data computation processing.

5) Subject knowledge acquisition methods: there are artificial, crawler or a combination of artificial crawler methods.

2.2.3. Overall process description. The process and technology mentioned in the previous two subsections are combined to complete the process design of discipline knowledge mapping, as shown in Figure 2.

The construction work of discipline mapping can be divided into three stages: preparation work stage, system construction stage, and display application stage. In the preparation stage, the overall construction work of the atlas is mainly planned according to the final use of the atlas. In the system construction stage, the core concepts of the discipline are mainly determined, and the overall hierarchical structure of the discipline concepts is determined according to the experts' opinions, the relationships between the knowledge points of the discipline are defined, and the resources corresponding to the knowledge entities and attributes are expanded. In the realization stage, the Protege software is used to complete the dataization of subject domain knowledge, and Neo4j is used to store and query the data to provide support for the subsequent secondary development.

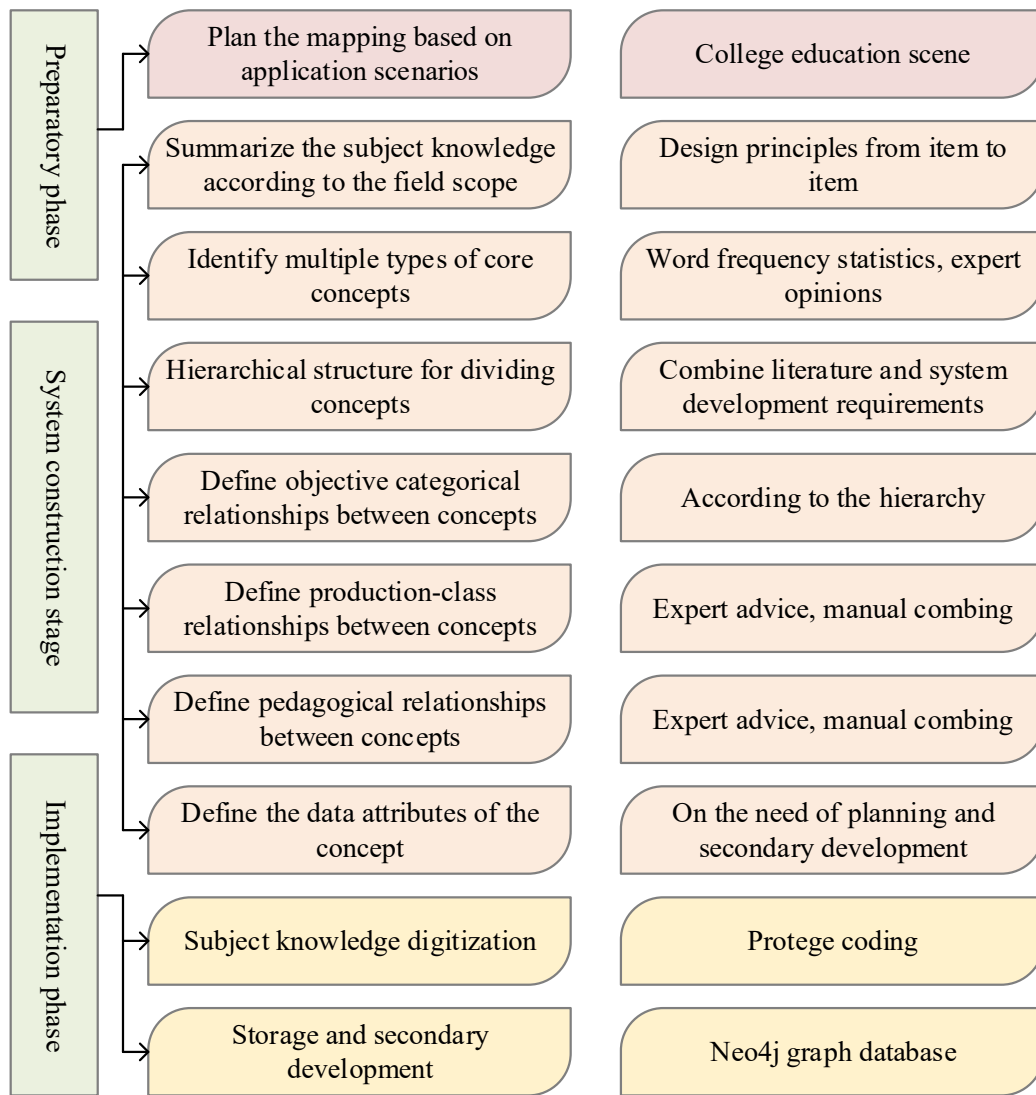


Fig. 2. Method design of map construction

3. Deep learning-based entity recognition and relationship extraction

3.1. Bidirectional long- and short-term memory networks

Bidirectional Long Short-Term Memory Network BiLSTM is a commonly used network model for feature encoding layer. BiLSTM consists of a forward and a reverse Long Short-Term Memory Network LSTM, so a brief introduction of LSTM is given before understanding BiLSTM.

The cell structure of LSTM contains three gates in each cell: current cell state and forgetting gate, input gate, and output gate.

The forgetting gate controls the forgetting of unwanted parts of past information and is computed as shown in Eq. where the upper unit output h_{t-1} is entered into the Sigmoid function with the current input x_t - the same as in the Sigmoid function - to give the current probability of forgetting f_t .

$$f_t = \sigma(W_f[h_{t-1}, x_t] + b_f) \quad (1)$$

The input gate calculation process is shown in Eq., similar to the process of the forgetting gate, the probability of the current input information is calculated first i_t , and then the tanh function is used to calculate the candidate cell state $\tilde{C}_t$, up to here there is the current forgetting probability f_t , the current input probability i_t , the cell state of the previous cell C_{t-1} , and the candidate cell state of the

current cell $\bar{C}_i$ needed for calculating the current cell state, and the cell of the current cell is calculated according to these parameters state C_i .

$$i_t = \sigma(W_i[h_{t-1}, x_t] + b_i) \quad (2)$$

$$\tilde{C}_t = \tanh(W_c[h_{t-1}, x_t] + b_c) \quad (3)$$

$$C_t = f_t * C_{t-1} + i_t * \bar{C}_t \quad (4)$$

The calculation process of the output gate is shown in Eq. vs. First, the output probability o_t is calculated, and then the output h_t of this cell is calculated based on the state of this cell C_i [33].

$$o_t = \sigma(W_o[h_{t-1}, x_t] + b_o) \quad (5)$$

$$h_t = o_t * \tanh(C_t) \quad (6)$$

BiLSTM is commonly used to further encode text vectors in the feature coding layer as its bidirectional network structure allows it to obtain semantic information and inter-word dependencies at a distance in the context compared to LSTM.

3.2. Conditional Random Fields

Conditional Random Field (CRF) is a commonly used classifier in entity recognition. Compared with other classifiers, CRF is able to learn the relationship between neighboring characters in the text, and avoid common sense errors when making predictions about character labels. For example, when using the BIO tagging method to represent an entity, B corresponds to the first character of the entity, I corresponds to the characters of the entity except the first character, and O corresponds to the non-entity characters, assuming that "BII" in the tagging sequence "OOBIIIOO" represents an entity, the prediction result of CRF after learning will not appear in the prediction result. Assuming that "BII" in the labeling sequence "OOBIIIOO" represents an entity, CRF will not make any obvious errors such as "IBI" or "OII" in the prediction result after learning.

CRF is a discriminative model that models the conditional probability distribution between the corresponding output sequences of the input sequences. CRF mainly consists of two parts, the firing matrix and the transfer matrix, the firing matrix represents the probability that each character in the text corresponds to each label, and the transfer matrix represents the probability that each label is transferred to other labels. CRF computes the maximum score corresponding to the text based on these two matrices, at which time the predicted tag sequence is the closest to the actual tag sequence of the text. The conditional probability calculation model is shown in Eq. The score calculation is shown in Eq. X denotes the input sequence $x_1, x_2, x_3, \dots, x_n$, Y denotes the output sequence $y_1, y_2, y_3, \dots, y_n$, P_{i,y_i} denotes the probability that the i th word belongs to tag y_i , W_{y_i, x_n} denotes the probability that tag y_i goes to tag y_i , and finally the total score is calculated.

$$P(Y | X) = \frac{\exp(\text{score}(x, y))}{\sum_y \exp(\text{score}(x, y'))} \quad (7)$$

$$\text{score}(X, Y) = \sum_{i=1}^n P_{i,y_i} + \sum_{i=1}^n W_{y_i, y_{i+1}} \quad (8)$$

3.3. Entity Recognition Based on BiLSTM+CRF Modeling

3.3.1. BiLSTM+CRF modeling. The advantage of LSTM is that it can learn the dependency between observation sequences through the bidirectional setting, and during the training process, LSTM can automatically extract the features of the observation sequences according to the target, but the

disadvantage is that it cannot learn the relationship between the state sequences, knowing that there is a certain relationship between the annotations in the named entity recognition task, for example, a class B annotation will not be followed by another class B annotation, so LSTM has the disadvantage of not being able to learn the context of the annotations, although it can save the very complicated feature engineering when solving NER such sequence annotation tasks.

On the contrary, CRF has the advantage of being able to model implicit states and learn the characteristics of state sequences, but it has the disadvantage of requiring manual extraction of sequence features [34]. So the general practice is to add a layer of CRF behind LSTM to obtain the advantages of both.

1) Feature Representation. In this method, the feature representation of words and phrases in the text is mainly performed using bidirectional LSTM, which is mainly designed to solve the problems encountered in the training process of traditional RNN models, such as gradient explosion and gradient vanishing.

Typically, an input sequence $(x_1, x_2, \dots, x_i, \dots, x_n)$ can be represented as $(\vec{h}_1, \vec{h}_2, \dots, \vec{h}_i, \dots, \vec{h}_n)$ using a forward LSTM, however, $\vec{h}_i$ can only represent all the information on the left side of x_i , ignoring the information on the right side, so another inverse LSTM needs to be utilized to model the information on the right side, which in turn can be used to represent the input sequence $(x_1, x_2, \dots, x_i, \dots, x_n)$ as $(\bar{h}_1, \bar{h}_2, \dots, \bar{h}_i, \dots, \bar{h}_n)$. In the end, the splicing of $\vec{h}_i$ and $\bar{h}_i$ is used h_i as the final representation of x_i .

2) Model training. After the above feature representation, many methods directly utilize the spliced vectors after bi-directional LSTM coding as the feature representation h_i for softmax classification, which in turn obtains the labels of each word, however, this method ignores the relationship between the labels, e.g., O can not be directly assigned to the label I after it, etc. Therefore the K-dimensional vectors (K is the number of labels) obtained from the representation of each word can be spliced to obtain the input P, which is a $n \times k$ -dimensional matrix that is uniformly input into the CRF model as features. The score of each sentence can be modeled as follows:

$$s(X, Y) = \sum_{i=0}^n A_{y_i, y_j} + \sum_{i=0}^n P_{i, y_i} \quad (9)$$

Where A is a transfer matrix, A_{y_i, y_j} represents the probability that a y_i label on the labeled sequence is transferred to the next y_j label. When the model is trained, the objective function is defined according to the score of each sentence mentioned above, and then training methods such as stochastic gradient descent are used to optimize the model parameters, and then the parameters of the whole network are trained, and in the process of training, in order to prevent overfitting, update rules such as Adadelta can be used.

3) Model classification. The classification stage is mainly to use the trained neural network model to classify the text to be classified. Firstly, BiLSTM is used to represent the features of the input text, and then it is input into the CRF to classify each word in the sentence, score the whole, and finally output the classification results to complete entity recognition.

3.3.2. Data pre-processing. Firstly, we collect the knowledge points of the subject-related content, and then use python crawler technology to crawl the content corresponding to the knowledge points from the relevant word information in Baidu encyclopedia, which provides the corresponding data base for the construction of knowledge graph. The specific crawling process is shown in Figure 3:

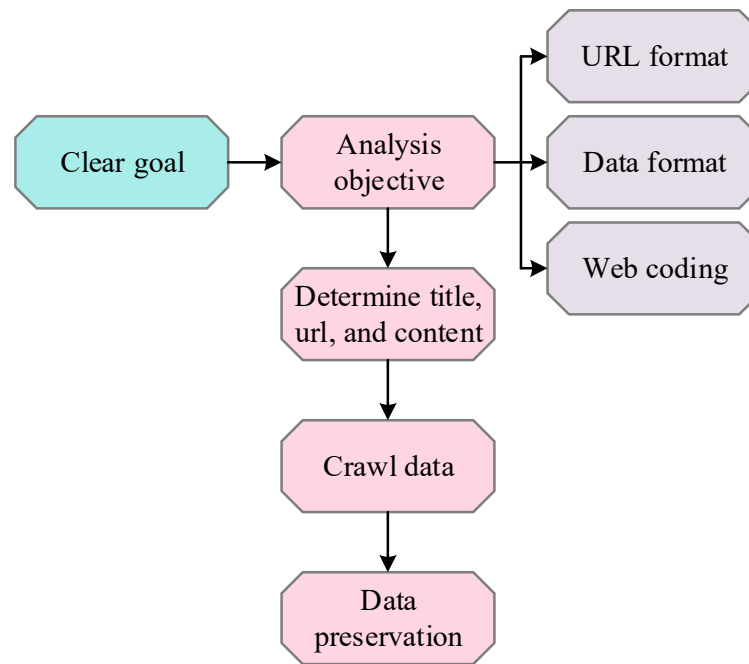


Fig. 3. crawler frame

1) Word Characteristics. Since Chinese does not have obvious word markers like English, so first of all, we have to divide it into words, there are many ways to divide words, but through comparison, this paper adopts the technology of jieba word division to divide it. jieba word division mainly has three modes: precise mode, full mode and search engine mode. However, due to the specificity of the subject of mathematics, for example, the mathematical term 'parallelogram', it can also be divided into 'parallel' and 'quadrilateral' which are two mathematical terms, thus In order to have more accurate effect of participle, this paper chooses the full mode of jieba participle for participle, so that the mathematical nouns divided are more comprehensive.

2) TF-IDF. TF-IDF is based on statistical features for entity recognition, which is used to assess the importance of a word to a document set or one of the documents in a corpus. The importance of a word increases proportionally with the number of times it appears in a document, but at the same time decreases inversely with the frequency of its appearance in a corpus. TF in TF-IDF stands for Word Frequency, which is a measure of how often a word occurs in a text, and IDF is the Inverse Document Frequency, which measures the general importance of the word in all document sets. The formula for TF-IDF is as follows:

$$W_{ij} = tf_{ij} \times idf_j = tf_{ij} \times \log \left(\frac{N}{n_j} \right) \quad (10)$$

where tf_i represents the frequency of occurrence of word t_j in document d_i , N represents the total number of documents in the corpus, and n_j represents the number of documents in which word t_j occurs.

The knowledge dataset has been obtained previously, which contains all the entities involved in this dataset, i.e., knowledge points. In this paper, we use the method of sequence annotation, the sequence annotation methods are mainly two kinds of BIO mode and BIOES, due to the better effect of BOIES annotation mode, so we use BIOES annotation mode for annotation.

3.4. Entity Relationship Extraction Based on BiLSTM+Attention Modeling

3.4.1. BiLSTM neural network model based on Attention mechanism. The BiLSTM+Attention model mainly consists of five hierarchies, i.e., Input Layer, Embedding Layer, LSTM Layer, Attention Layer and Output Layer. The first of these layers is the input layer, which mainly inputs sentences into the model; the second layer is the Embedding layer, which refers to mapping each word to a low-dimensional space; the third layer is the LSTM layer, which uses BiLSTM to obtain high-level features from the Embedding layer; and the fourth layer is the Attention layer, which generates a weight vector, and by multiplying with this weight vector multiplication, so that the word-level features in each iteration are merged into sentence-level features; the fifth level is the Output layer, which uses its own feature vector for relationship classification.

1) Word vector layer (Embedding layer). For a given sentence $S: S = x_1, x_2, \dots, x_T$ containing T word, each word x_i is converted to a vector e_i . For each word in S , there exists first a matrix of word vectors: $W^{wvd} \in R^{d^w \times |V|}$, V is a fixed-size vocabulary list, d^w is the dimensions of the word vectors, which are user-defined hyper-parameters, and W^{wrd} is a matrix of parameters learned through training. With this word vector matrix, each word is transformed into its word vector representation: $e_i = W^{wrd} v^i$, v^i are one-hot vectors of size $|V|$, with 1 at e_i and 0 at the other positions. Thus, sentence S will be transformed into a real matrix and passed to the next layer of the model.

2) Attention Layer. The set of vectors input to the LSTM layer is represented as $H: [h_1, h_2, \dots, h_T]$. The weight matrix obtained by the Attention layer consists of the following equations:

$$M = \tanh(H) \quad (11)$$

$$\alpha = \text{soft max}(w^T M) \quad (12)$$

$$r = H\alpha^T \quad (13)$$

where $H \in R^{d^w \times T}$, d^w are the dimensions of the word vectors and w^T is the transpose of the parameter vectors used for training and learning.

The final sentence used for classification is represented by the following equation:

$$h^* = \tanh(r) \quad (14)$$

3) Classification. In the classification task, a softmax classifier is used to predict the label $\hat{y}$. This classifier takes as input the obtained hidden state:

$$\hat{y}(y|S) = \text{soft max}(W^{(S)} h^* + b^S) \quad (15)$$

$$\hat{y} = \arg \max_y \hat{p}(y|S) \quad (16)$$

The cost function uses the negative log-likelihood of a positive sample:

$$J(\theta) = -\frac{1}{m} \sum_{i=1}^m t_i \log(y_i) + \lambda \|\theta\|_F^2 \quad (17)$$

where $t \in R^m$ is the one-hot representation of the positive sample, $y \in R^m$ is the probability of each category estimated by softmax (m is the number of categories), and λ is the hyperparameters of the $L2$ regularization.

3.4.2. Data annotation. The dataset of inter-entity relationships used in this study is formatted as "Entity 1 Entity 2 Inter-Entity Relationship Sentence", which starts with two entities, then the relationship between these two entities, and finally the sentence in which the two entities are located, separated by a space.

The dataset will be made by following all the sentences in the above manner and after that the work

of processing the dataset is to be done. While working on the dataset, first of all, following the words, the sentences have to be converted into word vectors. This requires us to create the word2id dictionary first and each word is converted to id. Then record the distance from each word in the sentence to each of the two entities. This gives us one word vector and two distance vectors for each word. The input to the model is the result of combining these three vectors. The output of the model is thus the id of the relationship between the knowledge points.

4. Knowledge integration

The knowledge fusion module has two tasks, the first one is to fuse the ternary data obtained from information extraction, mainly associating entities with the same name, which is achieved by entity disambiguation; the other task is to merge the structured data into the knowledge graph.

4.1. Entity Disambiguation

The ternary needed to construct the knowledge graph has been obtained in the information extraction, but since it is converted from a heterogeneous data source, although the authenticity of the data is guaranteed there will still be cases where the same entity exists with different names. This is when knowledge fusion is needed. The overall task of knowledge fusion is to calculate the similarity between entities, the similarity of entities within a certain threshold to be classified as the same entity, this paper's approach is to establish a sameAs relationship between the individual representations of the same entity, in the case of a user query on an entity, but also on the entity that has a sameAs relationship to do the same query. The entity similarity calculation method used in this paper is described below.

For node pairs in the knowledge graph, the similarity can be calculated based on their mapping relationship (parent-child/brother relationship) with the nearby entities, which is known as structure-level similarity. There are three categories of similarity calculation methods:

$$sim_s(C_1, C_2) = \mu_p sim_{string}(SC_1, SC_2) \quad (18)$$

$$sim_B(C_1, C_2) = \mu_B sim_B(BC_1, BC_2) \quad (19)$$

$$sim_R(C_1, C_2) = \mu_R sim_R(RC_1, RC_2) \quad (20)$$

(18)~(20) are the calculation methods of parent, child and brother similarity, respectively, for node c_1, c_2 , its parent node is SC_1, SC_2 , child node is BC_1, BC_2 , and brother node is RC_1, RC_2 , respectively. μ_p, μ_B, μ_R are the similarity decay coefficients in the parent rule, the child rule, and the brother rule, respectively. After calculating the above three similarities, the fusion is performed by weighted average, and the final similarity between nodes is obtained as:

$$sim_{structure}(C_1, C_2) = \frac{\alpha sim_s(C_1, C_2) + \beta sim_B(C_1, C_2) + \gamma sim_R(C_1, C_2)}{\alpha + \beta + \gamma} \quad (21)$$

α, β, γ is the weighted average coefficient, due to various types of rules have different impacts on nodes, usually $\alpha > \beta > \gamma$.

In order to test the effect of knowledge fusion, after fusion using inter-node similarity with a threshold of 0.5, this paper randomly selects 1000 fused entity pairs, and evaluates them by calculating their semantic similarity based on their attributes and relationships, and the formula for semantic similarity is:

$$sim_w(C_1, C_2) = \sum_{\rho(\omega C_1, \omega C_2) \in S(C_1, C_2)} Max(idf(\omega C_1), idf(\omega C_2)) \quad (22)$$

$Wordsim(\omega C_1, \omega C_2)$ denotes the word similarity of the attributes shared by node C_1, C_2 , $idf(\omega C_1)$

is the inverse text frequency index of ωC_1 , and $idf(\omega C_2)$ is the inverse text frequency index of ωC_2 .

4.2. MySQL Data Merge

The D2R tool maps data from a relational database to a virtual RDF format, with mapping rules for tables as shown in the table below. The mapping rules are similar to the cell labeling rules of the table-oriented data flow combination model.

4.3. OwnThink knowledge consolidation

The knowledge in OwnThink has been stored in CSV triples, and the data format is very standardized after cleaning and filtering, allowing direct use. For entities and relationships in OwnThink, it is necessary to add a label for them to match the ternary form of this paper, which is accomplished in this paper by using the method of building a label correspondence set and matching labels for each entity relationship. According to the "Neo4j Authoritative Guide" Neo4j provides several ways to import CSV data, but because the amount of data is too large, the use of LOADCSV and other commands will lead to memory overflow, so this paper uses the neo4j-adminimport command to import the knowledge of OwnThink.

4.4. Storage and visualization of disciplinary knowledge graphs

In the process of entity extraction and relationship extraction, the knowledge corpus is transformed into triples in the form of entity-relationship-entity, and these triples are stored in a graph database, which is adopted in this chapter. Neo4j graph database is a high-performance graph database that adopts a graph model to represent data, supports large-scale data storage and querying of data, and provides a flexible data model and the query language Cypher to support knowledge graph construction and querying.

4.4.1. Neo4j graph database. Neo4j is a graphical model-based database management system that can efficiently store and query large-scale data and can easily handle relationships between data. It is a mainstream knowledge graph technology widely used in recommender systems, social network analysis, and other data-intensive applications. In this chapter, Neo4j-community-4.2.17 version is chosen to implement the storage of knowledge graphs for data science disciplines.

4.4.2. Storage of disciplinary knowledge maps. In order to import the entity-relationship triples of knowledge graphs of data science disciplines into the Neo4j graph database, this chapter firstly categorizes these triples according to the module hierarchy, secondly saves them in different csv files according to the form of entity-relationship-entity, and then copies the sorted csv files into the import folder of the Neo4j graph database, and finally adopts the Cypher language to import all the csv files into Neo4j diagram database.

5. Simulation experiments and applications

5.1. Named Entity Recognition and Extraction Experiments

In this paper, we design an entity recognition based on BiLSTM+CRF model and an entity relationship extraction method based on BiLSTM+Attention model, respectively. In this section, experiments are designed to check their performance respectively.

5.1.1. Named Entity Recognition Experiment:

1) Experimental design. In this paper, the historical knowledge entity dataset constructed using manual annotation and semi-automatic annotation is used as the experimental data, which contains 8239 textual corpus, and the pre-processed dataset is divided into training set, testing set and

validation set according to 8:1:1. The Hidden Markov Model (HMM), Conditional Random Field (CRF), Long and Short-Term Memory Network (BiLSTM), BERT model and the BiLSTM-CRF model used in this paper are used to conduct comparison experiments on the self-constructed historical knowledge entity dataset, and the evaluation metrics are set to be Precision, Recall, and F1 value.

The CRF model used in this paper is a linear chain conditional random field, which is a typical discriminative model; the BERT model is the BERT-base-chinese proposed by Google; the BiLSTM and BiLSTM-CRF models use Word2Vec word embeddings based on Baidu Encyclopedia training, and Adam is selected as the optimizer of the model during the experiments, and the model specific The experimental parameters are shown in Table 1.

Table 1. Experimental parameter setting

Model parameter setting	Learning rate	Epoch	Dropout	Batch size	Max_seq_length
HMM	0.001	50	0.1	36	
CRF	0.001	50	0.1	36	
BiLSTM	0.001	50	0.1	36	
BERT	0.001	5	0.1	36	400
This model	0.001	50	0.1	36	

2) Analysis of experimental results. The experimental results are shown in Table 2, due to the Hidden Markov Model (HMM) adopts the chi-squared Markov assumption with observation independence assumption, resulting in poorer experimental effect compared with several other models; the linear chain conditional random field model is improved in effect by introducing customized feature functions to effectively overcome the problem of representation of dependency phenomenon in the Hidden Markov Model (HMM); due to the small size of the dataset, it is not enough for the CRF layer to learn more information between state transfers, so the BiLSTM-CRF model designed in this paper is improved compared with the BiLSTM model in terms of effect although the experimental effect is improved, but the improvement is small; the BERT model is based on the Transformer Encoder, which uses the mask language model to obtain the bidirectional fusion information, which is fine-tuned to be applied to downstream tasks.

Combining the experimental results with the research analysis in this paper, although the experimental effect of the BERT model is better than that of the BiLSTM-CRF model, considering that the BERT model consumes too much resources and takes a long time during training, and the difficulty of the task is small compared to the relationship extraction and attribute extraction, the gap between the experimental effect is not obvious, so the BiLSTM-CRF-based model designed in this paper is used for the subject teaching of the named entity recognition task is suitable.

Table 2. The named entity identifies the results of the experiment

Model	Accuracy rate	Recall rate	F1 score
HMM	0.8408	0.8539	0.8275
CRF	0.8194	0.8305	0.8109
BiLSTM	0.8343	0.8386	0.8324
BERT	0.9093	0.9272	0.9178
This model	0.8532	0.8841	0.8817

5.1.2. Named Entity Extraction Experiments:

1) Experimental design. The experiment takes the self-constructed historical knowledge relationship dataset as the experimental data, which contains 5537 textual corpus, and divides the pre-processed dataset into three parts, namely, training set, validation set and test set, according to 8:1:1. The evaluation indexes are set as Precision, Recall and F1 value. The Transformer model and PCNN model are selected for the experiments to compare with the model of BERT used in this paper on the self-constructed historical knowledge relationship dataset, in which the BERT model is the BERT-base-chinese model proposed by Google, and the Epoch is set to be 5, the Dropout to be 0.3, and the Batch size to be 32.

2) Analysis of experimental results. In this paper, PCNN model, Transformer model, BERT model and the model of this paper are chosen to carry out the task of historical knowledge attribute extraction in the experimental process, respectively. PCNN model, as a CNN network with simple structure, consists of the word vector of the word itself and the word's positional information relative to the entity word as the inputs, and then goes through convolutional layer, pooling layer and finally through Softmax function to score the output of each category. Scoring output for each category. Transformer encodes sentence information uniformly through encoder, extracts important features in the sentence through multi-head attention module, superimposed using residual network, and finally decodes using Decoder part of Transformer to get the relationship extraction result. BERT model is based on Transformer's Encoder part of Transformer as the encoder of the model, using the textual utterances as inputs to get a more comprehensive vector representation of the words when predicting each word, and inputting the obtained results to the fully connected layer for the relationship extraction task. The experimental results are shown in Table 3.

Compared to the Transformer model and BERT model, the model in this paper can better capture all the information in the input text sequence, and from the results in Table 3, it can be seen that the relationship extraction method based on the model in this paper outperforms the relationship extraction methods based on Transformer and BERT in terms of precision rate, recall and Macro-F1 score results on the self-constructed historical relationship dataset with 86.32%, 88.41%, and 88.17%, respectively.

Since the PCNN model only selects the most probable sentence from a certain entity pair of sentences for backpropagation calculation during the training process, there is an information loss problem. Compared with the PCNN model, the model in this paper is not only more convenient to use but also better than the PCNN model in the three evaluation indexes. Therefore, the combination of the model structure characteristics and the experimental results, the model in this paper is an efficient and excellent method for the relationship extraction task.

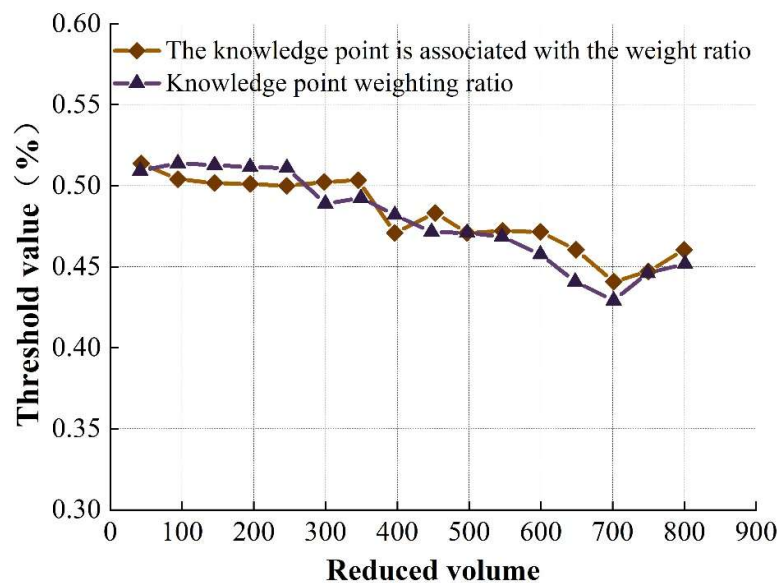
Table 3. Relationship extraction results

Model	Accuracy rate	Recall rate	F1 score
PCNN	0.8408	0.8539	0.8275
Transformer	0.8194	0.8305	0.8109
BERT	0.8343	0.8386	0.8324
This model	0.8632	0.8841	0.8817

5.2. Knowledge Subgraph Experiments

5.2.1. **Threshold Selection for Knowledge Subgraph Fusion.** The purpose of this experiment is to approximate the knowledge subgraphs in the subject knowledge graph, that is, to remove the redundant data in the knowledge subgraphs. Therefore, this experiment determines the relevant threshold value by selecting different numbers of knowledge subgraphs, which is the basis of knowledge subgraph fusion. In this regard, in order to determine the threshold selection problem in the process of fusion of subject knowledge subgraphs, this paper utilizes the ratio of subject knowledge point weight information and the ratio of knowledge point association weight information to design relevant experiments for determining the relevant thresholds, as shown in Figure 4.

It can be seen that the threshold value selected in this experiment is gradually reduced overall with the increase in the number of approximate diagrams. Therefore, the thresholds are negatively correlated for the set of knowledge subgraphs. From the above figure, it can be seen that when the number of subject knowledge subgraphs has reached 250, their knowledge point weights and knowledge point association weights are relatively unchanged. Therefore, this paper can be concluded that the results are: knowledge point weights and knowledge point association weights are between $[0.52, 0.54]$, $[0.48, 0.52]$, respectively. In this paper, the final threshold value is 0.52, which means that the knowledge association weight or knowledge point weight is less than the threshold value is deleted, and the remaining subgraphs are used for further calculation.

**Fig. 4.** Threshold selection experiment

5.2.2. **Similarity Assertion for Knowledge Subgraphs and its Performance Evaluation.** In order to validate the effectiveness of the proposed knowledge subgraph similarity calculation method in this

paper, it is compared with the mainstream methods such as Jaccard, DiceText, JaroDistance, and EditDistance. The evaluation criteria used in this experiment are Precision, F1 metric value. The experimental results are shown in Fig. 5 and Fig. 6. Among them, Fig. 5 shows the comparison of accuracy of different similarity methods and Fig. 6 shows the F1 value of different similarity methods.

Using the accuracy rate for comparison, it can be seen that under the same threshold, the method used in this paper has a relatively high accuracy in Precision compared to other methods, which shows the feasibility as well as the effectiveness of the method.

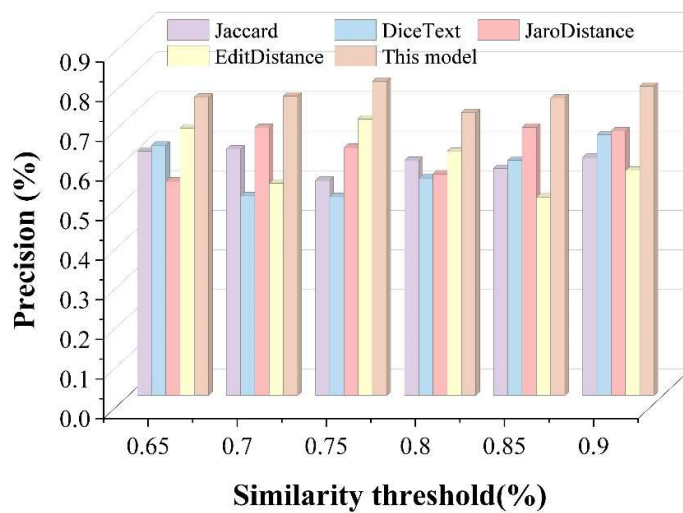


Fig. 5. Comparison of the accuracy method of different similarity methods

From the figure, it can be seen that this paper utilizes the metric F to compare the similarity calculation method proposed in this paper with the classical similarity calculation methods, and it can be concluded that the method proposed in this paper is higher than the other classical methods in terms of the F metric value, and what can be illustrated by this is the effectiveness of the method proposed in this paper.

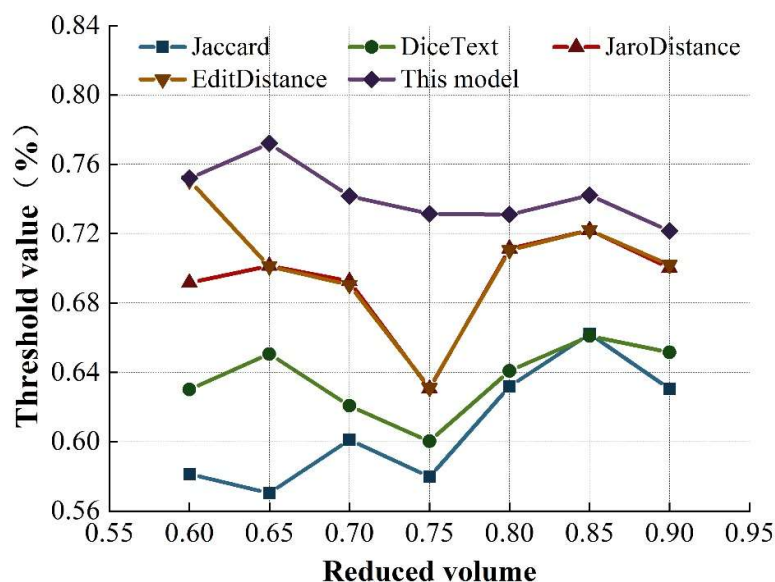


Fig. 6. F1 values of different similarity methods

5.3. Analysis of Knowledge Mapping Applications in Higher Education

A class is selected in a university to carry out the application experiment of the constructed knowledge graph, and the constructed knowledge graph is adopted for teaching. After conducting a one-month trial, the questionnaire survey method was used to collect and analyze the trial data to test the trial effect of the visualization system and the application effect of the knowledge graph in subject teaching. The questionnaire survey used whole group sampling, 150 copies were distributed and 126 copies were recovered. After eliminating invalid questionnaires, 120 copies were valid. Among them, 58 were boys and 62 were girls. SPSS software was used to process the data.

Cronbach's alpha was 0.8479, and the reliability of the questionnaire was high. The questionnaire was counted and analyzed from overall to local. As a whole, the class as a whole was satisfied with the trial of subject knowledge mapping. From a local point of view, the dimensions basically achieve the expected hypothesized effects, but there are also subtle gaps between the dimensions, as follows:

1) Overall analysis. In the overall analysis, in order to discern the overall attitudinal tendency of the class, the scores were counted according to 5=strongly agree, 4=agree, 3=uncertain, 2=disagree, and 1=strongly disagree. The number of people in each dimension was summed and then multiplied by the corresponding score, and the overall attitude tendency was judged by the total score. The results after scoring are shown in Figure 7. The bar chart visually reveals that the class as a whole holds a satisfactory attitude, with the largest number of people holding an agreeing attitude, and the total percentage of agreeing and strongly agreeing is more than 72%. Therefore, the effect of the use of knowledge mapping is generally in line with the design expectations.

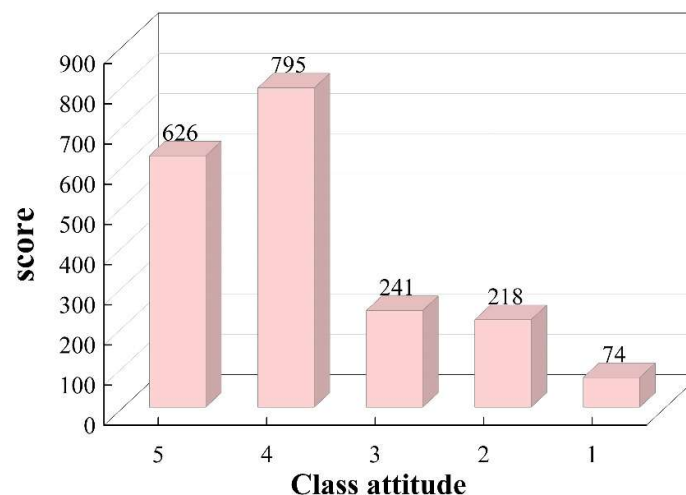


Fig. 7. General condition

2) System Design. Questions 1-3 of the questionnaire are the basic personal information part, and questions 4-6 are the specifics of the system design. The system design part is mainly investigated from three aspects: interface, operation and interaction, and the results of the survey are shown in Figure 8. In this part, the two attitudinal tendencies of agreeing and disagreeing account for a relatively high percentage, which is analyzed for the following reasons. First, at the initial stage of the construction of the visualization system, only part of the functions such as personal center, search mapping, and the structure of high school information technology courses were opened, thus affecting some students' experience of the system. Second, after the construction of the system was completed, students re-tried each function. Third, in the absence of supervision, some students did not experience

all the functions of the system as expected.

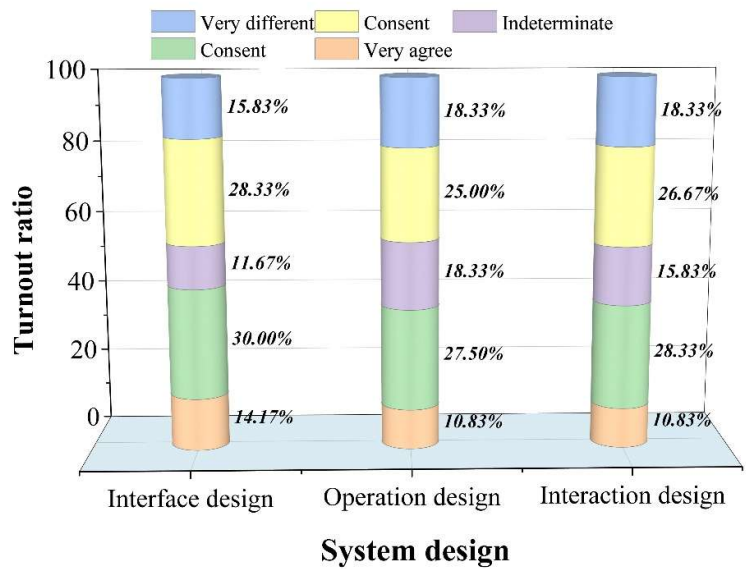


Fig. 8. System design

3) Learning support. Questions 7-11 of the questionnaire are specific to learning support, and this part is mainly investigated in three aspects: knowledge presentation, knowledge construction and teaching resources, and the results of the survey are shown in Figure 9.

Among them, question 7 is the knowledge presentation part, which mainly investigates whether the subject knowledge mapping integrates knowledge and clearly demonstrates the subject knowledge system. In this part, the total percentage of agree and strongly agree is 64.17%. The data indicate that most learners believe that the knowledge mapping visualization system is able to present the conceptual structure clearly, explicitly, and systematically. Questions 8-10 are the knowledge construction part, and the three headline items of questions 8-10 are summarized into a whole by calculating the variables when SPSS counts the knowledge construction dimensions. Question 11 is the teaching resources section, and this section investigates the impact of digital educational resources integrated in the subject knowledge mapping visualization system on knowledge retrieval. A total of 60.83% agreed and strongly agreed, which means that more than 60% of the learners believed that the digital teaching resources integrated in the knowledge graph visualization system improve knowledge retrieval.

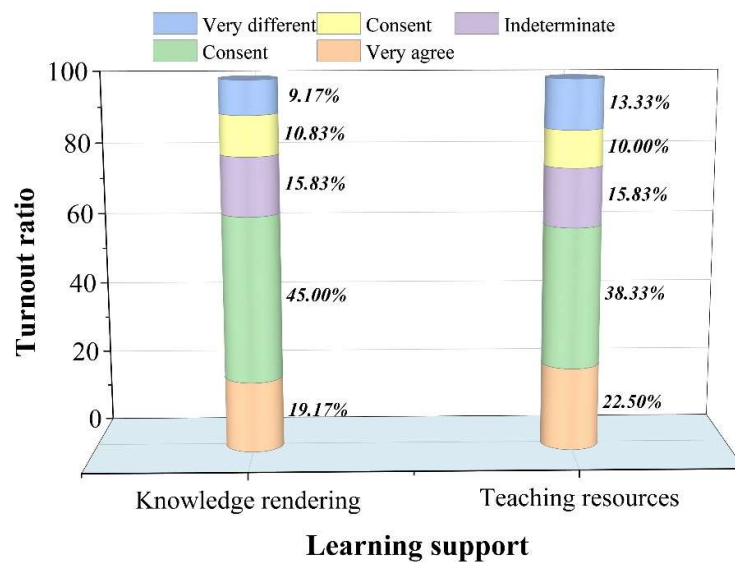


Fig. 9. Learning support

4) Knowledge construction. The data results of the knowledge construction part are shown in Figure 10. As a whole, more than 70% of the learners believe that knowledge mapping is conducive to the correlation of previously learned knowledge with what they are learning now, and to the fusion of old and new knowledge, and that it is easy for learners to form their own knowledge structure system under the directed guidance of knowledge mapping. In terms of the specific distribution, in questions 8, 9 and 10, the sum of the percentages of agree and strongly agree is 69.16%, 74.16% and 70.83% respectively. The data indicate that the largest number of people agree and strongly agree in question 9, i.e., learners believe that knowledge mapping has a facilitating effect on associative knowledge.

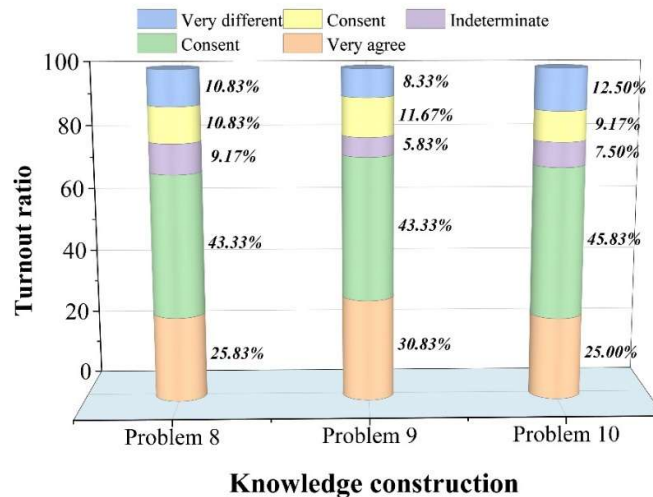


Fig. 10. Knowledge building

5) Learning Effectiveness. Questions 12-14 of the questionnaire are specific to learning effectiveness. Learning effect is mainly investigated from three aspects: learning efficiency, learning input and tool utilization, as shown in Figure 11. Learning efficiency mainly refers to the efficiency of learners in using knowledge mapping to extract existing knowledge. Among them, the total percentage of agree and strongly agree is 67.5%. Learning engagement mainly reflects whether knowledge mapping can attract learners' attention and focus their energy, with a total of 64.16% agreeing and strongly agreeing.

Tool use, on the other hand, refers to the learners' migratory use of learning tools, with a total of 61.67% agreeing and strongly agreeing.

Overall, among the three dimensions, learners hold the most agreement and strong agreement on learning support, and the knowledge construction component has reached more than 70%. The Learning Effectiveness dimension is slightly next in importance, with both the percentage of Agree and Strongly Agree and both ranging from 60% to 70%. The system design component, on the other hand, shows a bifurcated result of agreement and disagreement. On the one hand, the data presents the trial effect of the knowledge graph visualization system, and on the other hand, it reflects the gap between the dimensions and the deficiencies in the construction of the visualization system, which has been optimized appropriately for the deficiencies in the design of the system in this study.

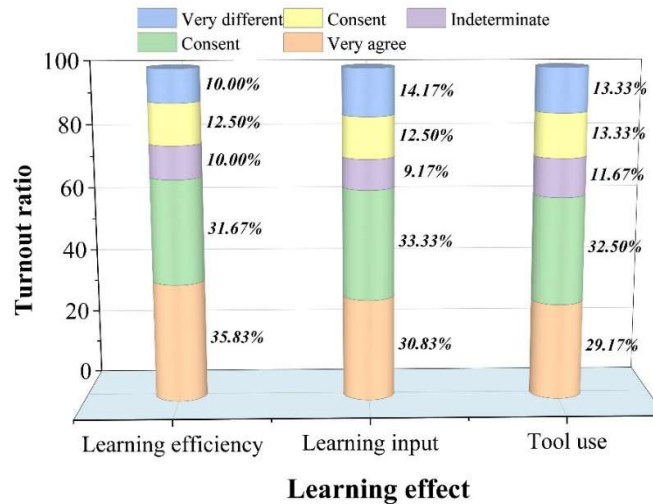


Fig. 11. Learning effect

6. Conclusion

In order to follow the trend of education informatization, a BiLSTM-based knowledge graph of college subjects is constructed. Its performance is verified on a dataset and applied to teaching in a university. The research results are as follows:

- 1) The BiLSTM-CRF model designed in this paper has a small improvement compared with the BiLSTM model although the experimental effect is improved. Considering that the BERT model consumes too much resources and takes a long time during training, the BiLSTM-CRF-based model designed in this paper is appropriate for the task of named entity recognition for teaching subjects.
- 2) Compared with the Transformer model and the BERT model, the model in this paper can better capture all the information in the input text sequences, and the results of its extraction method on the self-constructed historical relational dataset are 86.32%, 88.41%, and 88.17% in terms of precision, recall, and Macro-F1 score, respectively, which are better than those of the Transformer-based and the BERT-based relationship extraction methods.

Funding

SSelf-Funded Project for Science and Technology Research and Development Plan of Langfang City in 2024 -- Construction and Application of University Disciplinary Knowledge Graph Driven by Educational Big Data (Project Number: 2024011027).

References

- [1] Gomez-Perez, J. M., Pan, J. Z., Vetere, G., & Wu, H. (2017). Enterprise knowledge graph: An introduction. In *Exploiting linked data and knowledge graphs in large organisations* (pp. 1-14). Cham: Springer International Publishing.
- [2] Wu, W., Wen, C., Yuan, Q., Chen, Q., & Cao, Y. (2023). Construction and application of knowledge graph for construction accidents based on deep learning. *Engineering, construction and architectural management*.
- [3] Abu-Salih, B., & Alotaibi, S. (2024). A systematic literature review of knowledge graph construction and application in education. *Heliyon*.
- [4] Wu, X., Wu, J., Fu, X., Li, J., Zhou, P., & Jiang, X. (2019, November). Automatic knowledge graph construction: A report on the 2019 icdm/icbk contest. In *2019 IEEE International Conference on Data Mining (ICDM)* (pp. 1540-1545). IEEE.
- [5] Vakaj, E., Tiwari, S., Mihindukulasooriya, N., Ortiz-Rodríguez, F., & Mcgranaghan, R. (2023, April). NLP4KGC: natural language processing for knowledge graph construction. In *Companion Proceedings of the ACM Web Conference 2023* (pp. 1111-1111).
- [6] Fang, W., Ma, L., Love, P. E., Luo, H., Ding, L., & Zhou, A. O. (2020). Knowledge graph for identifying hazards on construction sites: Integrating computer vision with ontology. *Automation in Construction*, 119, 103310.
- [7] Yu, T., Li, J., Yu, Q., Tian, Y., Shun, X., Xu, L., ... & Gao, H. (2017). Knowledge graph for TCM health preservation: Design, construction, and applications. *Artificial intelligence in medicine*, 77, 48-52.
- [8] Jiang, T., Zhao, T., Qin, B., Liu, T., Chawla, N. V., & Jiang, M. (2019, July). The role of " condition" a novel scientific knowledge graph representation and construction model. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining* (pp. 1634-1642).
- [9] Zhu, Y., Wang, X., Chen, J., Qiao, S., Ou, Y., Yao, Y., ... & Zhang, N. (2024). Llms for knowledge graph construction and reasoning: Recent capabilities and future opportunities. *World Wide Web*, 27(5), 58.
- [10] Tan, J., Qiu, Q., Guo, W., & Li, T. (2021). Research on the construction of a knowledge graph and knowledge reasoning model in the field of urban traffic. *Sustainability*, 13(6), 3191.
- [11] Chen, Z., Wang, Y., Zhao, B., Cheng, J., Zhao, X., & Duan, Z. (2020). Knowledge graph completion: A review. *Ieee Access*, 8, 192435-192456.
- [12] Cardoso, S. D., Da Silveira, M., & Pruski, C. (2020). Construction and exploitation of an historical knowledge graph to deal with the evolution of ontologies. *Knowledge-Based Systems*, 194, 105508.
- [13] Li, L., Wang, P., Yan, J., Wang, Y., Li, S., Jiang, J., ... & Liu, Y. (2020). Real-world data medical knowledge graph: construction and applications. *Artificial intelligence in medicine*, 103, 101817.
- [14] Li, X., Zhang, S., Huang, R., Huang, B., & Wang, S. (2019). Process knowledge graph construction method for process reuse. *Xibei Gongye Daxue Xuebao/Journal of Northwestern Polytechnical University*, 37(6), 1174-1183.
- [15] Ko, H., Witherell, P., Lu, Y., Kim, S., & Rosen, D. W. (2021). Machine learning and knowledge graph based design rule construction for additive manufacturing. *Additive Manufacturing*, 37, 101620.
- [16] Rizun, M. (2019). Knowledge graph application in education: a literature review. *Acta Universitatis Lodziensis. Folia Oeconomica*, 3(342), 7-19.
- [17] Li, F. L., Chen, H., Xu, G., Qiu, T., Ji, F., Zhang, J., & Chen, H. (2020, October). AliMeKG: Domain knowledge graph construction and application in e-commerce. In *Proceedings of the 29th ACM International Conference on Information & Knowledge Management* (pp. 2581-2588).

- [18] Schroeder, N. L., Nesbit, J. C., Anguiano, C. J., & Adesope, O. O. (2018). Studying and constructing concept maps: A meta-analysis. *Educational psychology review*, 30, 431-455.
- [19] Zhu, Y., Chen, D., Zhou, C., Lu, L., & Duan, X. (2021, July). A knowledge graph based construction method for Digital Twin Network. In *2021 IEEE 1st International Conference on Digital Twins and Parallel Intelligence (DTPI)* (pp. 362-365). IEEE.
- [20] Wette, R. (2017). Using mind maps to reveal and develop genre knowledge in a graduate writing course. *Journal of second language writing*, 38, 58-71.
- [21] Lou, P., Yu, D., Jiang, X., Hu, J., Zeng, Y., & Fan, C. (2023). Knowledge Graph Construction Based on a Joint Model for Equipment Maintenance. *Mathematics*, 11(17), 3748.
- [22] Molinari, G. (2017). From learners' concept maps of their similar or complementary prior knowledge to collaborative concept map: Dual eye-tracking and concept map analyses. *Psychologie française*, 62(3), 293-311.
- [23] Haussmann, S., Seneviratne, O., Chen, Y., Ne'eman, Y., Codella, J., Chen, C. H., ... & Zaki, M. J. (2019). FoodKG: a semantics-driven knowledge graph for food recommendation. In *The Semantic Web-ISWC 2019: 18th International Semantic Web Conference, Auckland, New Zealand, October 26-30, 2019, Proceedings, Part II 18* (pp. 146-162). Springer International Publishing.
- [24] Chen, P., Lu, Y., Zheng, V. W., Chen, X., & Yang, B. (2018). Knowedu: A system to construct knowledge graph for education. *Ieee Access*, 6, 31553-31563.
- [25] Zhang, N., Zhang, W., & Shang, Y. (2021, November). Research on Dynamic Knowledge Map Service System Using Computer Big Data. In *Journal of Physics: Conference Series* (Vol. 2083, No. 4, p. 042001). IOP Publishing.
- [26] Melnyk, I., Dognin, P., & Das, P. (2022, December). Knowledge Graph Generation From Text. In *Findings of the Association for Computational Linguistics: EMNLP 2022* (pp. 1610-1622).
- [27] Deng, Y. (2019). Construction of higher education knowledge map in university libraries based on MOOC. *The Electronic Library*, 37(5), 811-829.
- [28] Yang, H., Xin, R., Zhang, N., Sun, C., & Pan, J. (2023, September). Research on Constructing System of Graphic Knowledge Map Based on Multi-Modal Data. In *2023 IEEE 6th International Conference on Information Systems and Computer Aided Education (ICISCAE)* (pp. 855-860). IEEE.
- [29] Nitchot, A., Wettayaprasit, W., & Gilbert, L. (2017, March). Monitoring, teaching and learning using knowledge maps and structures (MyTeLeMap). In *2017 International Conference on Digital Arts, Media and Technology (ICDAMT)* (pp. 189-194). IEEE.
- [30] Zou, X., Yue, W. L., & Le Vu, H. (2018). Visualization and analysis of mapping knowledge domain of road safety studies. *Accident Analysis & Prevention*, 118, 131-145.
- [31] Paulheim, H. (2017). Knowledge graph refinement: A survey of approaches and evaluation methods. *Semantic web*, 8(3), 489-508.
- [32] Jingni Song, Luyi Bai, Xuanxuan An & Longlong Zhou. (2025). Unsupervised fuzzy temporal knowledge graph entity alignment via joint fuzzy semantics learning and global structure learning. *Neurocomputing* 129019-129019.
- [33] Zixi Chen, Matteo Bernabei, Vanessa Mainardi, Xuyang Ren, Gastone Ciuti & Cesare Stefanini. (2024). A Novel and Accurate BiLSTM Configuration Controller for Modular Soft Robots with Module Number Adaptability.. *Soft robotics*
- [34] Kaushik Bose & Kamal Sarkar. (2024). Named Entity Recognition in Bengali and Hindi Using MuRIL and Conditional Random Fields. *SN Computer Science*(7), 856-856.