

A Personalized Path Generation Model for English Learning Based on Natural Language Processing

Danbai Liu^{1, ✉}, Yongning Qian², Jing Zhao³

¹ The School of Humanities and Arts, Shaanxi Technical College of Finance and Economics, Xianyang, Shaanxi, 712000, China

² The School of Humanities and Arts, Shaanxi Technical College of Finance and Economics, Xianyang, Shaanxi, 712000, China

³ The School of Accounting, Shaanxi Technical College of Finance and Economics, Xianyang, Shaanxi, 712000, China

ABSTRACT

In this paper, based on the knowledge graph, word vectors and other personalized path generation related technologies, based on the graph convolutional neural network to complete the construction of the English knowledge graph model, to generate a personalized English knowledge graph, drawing on the data structure in the graph, to generate a personalized learning path, in order to make the generation of personalized learning path is more reasonable, in accordance with the difficulty value of the exercises for the exercises to be sorted. Simulation experiments are designed to evaluate the difficulty level of the generated exercises. The difficulty level of most of the English exercises generated by the personalized recommendation path is concentrated in the easy and general levels, and there are a total of 2,229 questions in these two difficulty levels, so the difficulty level of the generated questions is moderate. After a period of personalized path-generated English learning, six teaching activities were carried out, and the average score of the first post-test of the experimental group was higher than that of the control group, and the Sig values were all less than 0.05, indicating that the difference in the scores of the two groups of students was significant, which side by side reflected the accuracy of personalized path-generated English teaching.

Keywords: graph convolutional neural network, word vector, English knowledge map, personalized recommendation path, English learning

1. Introduction

Natural Language Processing (NLP) is an important technology in the field of Artificial Intelligence, aiming to enable machines to understand, interpret and generate natural language. The goal of NLP is to enable computers to understand and process various forms of natural language as humans do, and it is widely used in the fields of text, speech and image processing [1-4]. In text analysis, NLP can be

✉ Corresponding author.

E-mail address: 15991090109@163.com (D. Liu).

Received 15 February 2024; Revised 05 June 2024; Accepted 25 October 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-061](https://doi.org/10.61091/jcmcc127b-061)

© 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/>).

used to extract key information from a large amount of text data, perform sentiment analysis, text classification and so on. In the field of machine translation, NLP can realize the automatic translation function to translate one language into another language [5-8]. In the development of intelligent assistants and virtual assistants, NLP can achieve speech recognition and speech synthesis, allowing users to interact with devices through voice, which has important applications in the generation of personalized paths for English learning [9-11].

With the continuous development and popularization of educational technology, personalized learning has become a hot topic in the field of education. There is an obvious problem in the traditional English teaching model, that is, it does not take into account the different needs and learning styles of each student. And the optimization of personalized learning path design and learning support system is to solve this problem [12-15]. Personalized learning path design refers to designing a unique learning path for students according to their individual differences and learning characteristics to meet their individual learning needs. Compared with the traditional unified teaching mode, personalized learning path design has the important role of meeting students' individual needs, stimulating students' interest in learning, and enhancing learning effectiveness [16-19].

Literature [20] describes a data-driven personalized learning model aimed at improving foreign language acquisition. Experiments on two educational platforms, Edmodo and Duolingo, show that the model exhibits potential effectiveness in educational settings and facilitates personalized education. Literature [21] introduces the RPR algorithm designed to facilitate personalized English listening learning. A series of experiments and analyses showed that the algorithm was effective in improving student engagement and comprehension by organizing listening materials, contributing to a personalized and effective language learning experience. Literature [22] carried out the design and evaluation of a personalized English learning system based on artificial intelligence, and based on the analysis of user satisfaction surveys, technical performance, and other indicators, it showed that the advantages of the system in supporting personalized learning needs and other aspects are obvious. Literature [23] evaluates students' vocabulary level based on natural language processing methods. And combined with reinforcement learning algorithms, it realizes the dynamic adjustment of the learning process. The experimental results emphasize that reinforcement learning is conducive to improving students' engagement. Literature [24] explored the application of intelligent educational technology in business English education. Using an interpretive sequential mixed methods research design and based on a quantitative survey of students showed that intelligent educational technology helps to improve students' business English competence and supports the development of personalized instruction. Literature [25] explores the role of AI in education, particularly personalized learning pathways for language education. It also advocates a balanced approach that utilizes the advantages of AI while retaining essential human factors and inclusiveness, thus contributing to the development of a digital education system.

Literature [26] describes the application of Natural Language Processing (NLP) in language teaching. It points out that NLP has advantages such as enhanced personalization and future directions such as developing interaction and personalization. It also reveals its shortcomings such as data privacy. Literature [27] examined the application of AI in MALL and its impact on English language teaching. Based on the literature analysis of Web of Science database, it was found that AI has the advantages of realizing personalized learning content and improving the effect of English learning, which helps to improve the quality of teaching. Literature [28] explores the potential of AI language learning tools including technologies such as NLP in reshaping traditional language education. It emphasizes that AI increases student engagement and can improve learning outcomes. And reveals the challenges it poses. Literature [29] explores the use of AI in personalized language learning. Based on the literature review, the aim was to assess the effectiveness, challenges and potential of AI tutoring systems. The findings suggest that AI presents opportunities for language

education, as well as challenges such as privacy and security. Literature [30] proposes techniques based on AI-L-TM to recommend learning paths with IoT devices, data mining methods, etc., aiming to improve the English learning experience of university students. The advantages and disadvantages of these techniques and their importance in creating smarter online learning environments are discussed. Literature [31] investigates the effectiveness of AI tools and methods in language learning, especially in improving learning outcomes. It reveals the important role of AI in providing instant feedback, immersive learning in several aspects, and realizing an efficient language learning experience.

In this paper, we apply word2vec word vector model, combined with jieba split word tool to perform split word operation on parsing, and pass the parsing and the English vocabulary involved in it to the word2vec model for training, to get each English vocabulary generating word vectors, and parse to get sentence vectors in the same way. Using graph convolutional neural network, the English knowledge map neighbor matrix is generated, and the RELU function is chosen as the activation function to improve the construction of English knowledge map. After preprocessing, cleaning, extracting, fusing and labeling the data from the data source, the knowledge graph and learner user profiles are finally established to generate the English personalized learning recommendation path. The exercises that make up the personalized learning path are selected in the personalized knowledge graph through the students' doing and the similarity between the exercise nodes. After that, personalized learning paths are generated according to the difficulty level of the exercises. Simulation experiments are set up to evaluate the difficulty level of the exercises. At the same time, combined with empirical analysis, the English learning effect based on personalized path generation is evaluated.

2. Personalized Path Generation for English Learning Based on Natural Language Processing

2.1. *Personalized path generation related technology*

2.1.1. Knowledge mapping. The data layer of knowledge graph is stored in the graph database with "entity-attribute-value" ternary as the expression of facts [32]. The ontology knowledge base is the abstract conceptual framework of knowledge mapping, which can be used to abstractly categorize the related knowledge when it comes to university English, for example, English teaching can be divided into listening, speaking, reading, writing and translation. The underlying database saves entity relationships and entity attribute values schema layer is built on top of the data layer, which is the core of the knowledge graph, and the knowledge stored in the schema layer is the refined knowledge, and ontology libraries are usually used to manage the schema layer of the knowledge graph, and the ontology libraries are used to standardize the connection between entities, relationships, and entity types and attributes with the help of the support of the ontology libraries for the axioms, rules, and constraints.

The whole process relies on natural language processing, relationship extraction, knowledge representation, machine learning and other techniques, and finally constructs a node-rich and relationship-rich knowledge graph. The knowledge graph encodes domain-specific knowledge elements and their associations and structures, providing the basis for the next step of knowledge reasoning, question and answer, learning and other applications. The knowledge graph provides an information base for the optimization of the teaching model. By analyzing the knowledge structure and learners' data in the knowledge map, the learners' knowledge deficiencies and needs can be found, which helps to adjust the teaching progress and focus and realize accurate teaching. At the same time, the knowledge graph can also recommend personalized learning paths according to the knowledge status of the learners and guide them to learn independently. The continuous updating of knowledge in the knowledge graph helps to upgrade the teaching content in time. When the

knowledge graph is expanded by adding or improving new information, the teaching content needs to be adjusted accordingly to ensure the timeliness. This requires teachers to pay constant attention to the update of the knowledge graph and reflect it in teaching in a timely manner. Learning analysis and recommendation based on knowledge graph can realize the personalized transformation of learning style. Different learners will get different knowledge supplements or exercises, some focusing on listening, some focusing on speaking, etc., which can help learners personalize their learning guided by their individual interests and needs.

2.1.2. Word vectors:

1) Word2vec model. Word2vec model is a classic local context word vector model applied in the field of short text classification in the early days, its main idea is based on a two-layer neural network, the input is the One-hot encoding corresponding to the context of the target word in the text corpus, and the output is a set of numeric vectors to represent each word in the corpus, which contains two training models, CBOW and Skip-gram. These two models are similar to Neural Network Language Model (NNLM), the difference is that NNLM is to train the language model, and the word vectors are just as a by-product. The direct purpose of CBOW and Skip-gram models is to obtain high-quality word vectors, and simplify the training steps to optimize the synthesis method, which directly reduces the computational complexity. However, Word2vec model is a model based on static word vectors, due to its one-to-one vector representation, it can not solve the problem of multiple meanings of words well, and can not be dynamically optimized for specific tasks.

The Word2vec model formula is shown in (1), $p(t|c)$ is the probability of the target word t appearing before and after the context word c , e_c is the embedding of c , θ is the parameter for linear computation of the Softmax layer, which is a normalized network aiming at mapping multiple scalars into a probability distribution, and each of its outputs ranges in the range of (0, 1), and x is the vocabulary size:

$$p(t|c) = \frac{e^{\theta_i^T e_c}}{\sum_{j=1}^x e^{\theta_j^T e_c}} \quad (1)$$

The dimension of Word2vec word vectors is usually between 100-300 dimensions, each dimension represents the shallow semantic features of the word, and the similarity between words is reflected by the distance between the vectors, which makes the word vectors generated by the Word2vec model widely used in the field of short text categorization [33].

2) Glove model. Since Word2vec is trained by a local corpus, its feature extraction is based on sliding window, in order to improve the problem that the semantic information of Word2vec local corpus training is not rich enough. In order to utilize the lexical co-occurrence information as much as possible, Glove model constructs a lexical co-occurrence matrix, which can be trained on large-scale corpus faster and easier. Glove is compared with Word2vec model on the word category task. Glove significantly outperforms the comparative baseline model, indicating that Glove better captures the semantics and relationships between words and produces more accurate word vector representations. However, the Glove model uses global information, which is relatively memory intensive compared to the other models, and compared to Word2vec, the Glove still cannot completely solve the problem of word polysemy.

The formula of Glove model is shown in (2), X_{ij} represents the frequency of word i and word j appearing in the same window. v_i and v_j are the word vectors of words i and j , b_i and b_j are two scalars, N represents the size of the vocabulary list, $f(x)$ is a weight function, the higher the frequency the greater the weight, the formula is shown in (3):

$$J = \sum_{i,j}^N f(X_{i,j})(v_i^T v_j + b_i + b_j - \log(X_{i,j})) \quad (2)$$

$$f(x) = \begin{cases} \frac{x}{x_{\max}} & x < x_{\max} \\ 1 & x \geq x_{\max} \end{cases} \quad (3)$$

Glove model is an effective word vector representation model in short text categorization. Based on global word frequency statistics, it is able to learn a more accurate, faster and effective word vector representation. In the short text classification task, Glove model can help the model to better understand the semantics of the text.

2.1.3. Graph Volume Neural Networks. Graph neural network refers to the neural network that can deal with the graph data, before the existence of convolutional neural network and its variants and recurrent neural network and its variants can only deal with the data in the Euclidean domain, that is, two-dimensional spatial data, there is no better way for the non-Euclidean domain of the graph data structure, and the real-life graph data structure is common and important, through the introduction of the graph neural network can be very good to deal with these graph data, make full use of the graph data, solve the problem that can not be solved before [34].

2.1.4. Learning path generation. The minimum spanning tree in a weighted graph is the path that contains all the nodes in the graph with the smallest sum of weights, and there are several methods for generating the minimum spanning tree based on the weighted graph, the most classical being the kruskal algorithm and the prim algorithm. The Kruskal algorithm is commonly known as the "edge addition method", in the initial path, there are only nodes and no edges in the path, and each iteration selects an edge with the smallest weight to join the path, until the number of edges in the path is less than the number of nodes.1.The prim algorithm is commonly known as the "point addition method", there is only one node in the path at the beginning, and the point corresponding to the edge with the smallest weight is selected in the subsequent iteration to be added to the path, until the number of edges in the path is less than the number of nodes, through the above two methods, Based on the undirected weighted graph, it is easy to obtain the minimum spanning tree, and the graph path with a specified number of nodes can also be obtained [35].

2.2. English Knowledge Graph Construction

2.2.1. Data preparation. When generating the feature vectors corresponding to the exercise data in the test set, the feature vectors corresponding to the exercises in the training set, and the feature vectors corresponding to the labeled exercise data in the training set, we need to use word2vec to generate feature vectors for each exercise, and in this paper the dimension of the feature vectors is set to 512. The most important part of the exercise data is the stem and the parse, so we need to consider only the stem and the parse when generating feature vectors for the exercise data. The most important parts of the exercise data are the stem and the parse, so only the stem and the parse of the exercise need to be considered when generating feature vectors for the exercises. When generating feature vectors for the stem and parses of an exercise, it is necessary to firstly perform lexical segmentation on the stem and parses of the exercise to obtain all the words involved in the exercise data. The stem part of an exercise consists of English words separated by spaces, and all the words in the stem can be obtained by using spaces to slice the stem. Most of the parsing part of the exercise is in English, in the split word will need to use the English split word tool, there are many English split word tool, this paper chooses jieba split word tool for parsing split word operation. After obtaining all the words involved in the exercise data, the exercise stem and the English words involved in it are

passed to the word2vec model for training, and the word vector can be generated for each English word, and after obtaining the word vector corresponding to each word, the word vector of all the words contained in the exercise stem is added and divided by the number of words contained in the question stem n to obtain the sentence vector SQ corresponding to the exercise stem, as shown in equation (4). Where SQ_k denotes the value of the sentence vector in the k th dimension of the most SQ and W_{ik} denotes the value of the i th word vector in the stem in the k th dimension:

$$SQ = (SQ_1, SQ_2 \dots SQ_{512}) = \left(\frac{\sum_{i=1}^n W_{i1}}{n}, \frac{\sum_{i=1}^n W_{i2}}{n} \dots \frac{\sum_{i=1}^n W_{i512}}{n} \right) \quad (4)$$

Similarly, the parsing and the English vocabulary involved are passed to the word2vec model for training, then a word vector can be generated for each English vocabulary, and then the corresponding sentence vector of the parsing can be obtained in the same way, and after that, the sentence vector corresponding to the question and the sentence vector corresponding to the parsing can be summed up and divided by 2 to obtain the feature vector ST of the exercise as shown in Equation (5). Where ST_k denotes the value in dimension k of the feature vector, SQ_k denotes the value in dimension k of the sentence vector, and SA_k denotes the value in dimension k of the vector corresponding to the parsing of the exercise:

$$ST = (ST_1, ST_2 \dots ST_{512}) = \left(\frac{SQ_1 + SA_1}{2}, \frac{SQ_2 + SA_2}{2} \dots \frac{SQ_{512} + SA_{512}}{2} \right) \quad (5)$$

2.2.2. Graph Convolutional Neural Networks. Graph Convolutional Neural Network is a neural network model that performs convolution operation on graphs, this convolution operation has many similarities with the convolution operation in convolutional neural networks such as weight sharing, local connectivity, multilayer networks, etc. There are many exercise nodes and interconnected edges on the English knowledge graph, when representing one of the exercise nodes, the model will use its surrounding nodes to represent this node, which is the process of graph convolution. In the process of convolution, the features of the exercise nodes around the exercise node will have an impact on the representation of the current exercise node, the essence of this process is feature extraction, with the increasing number of times the model is trained, deeper features are extracted. Graph Convolutional Neural Networks have multiple hidden layers, and in multilayer graph convolutional neural networks, the propagation rule between layers is defined by equation (6):

$$H^{(l+1)} = \sigma \left(\check{D}^{-\frac{1}{2}} (A + I_N) \check{D}^{-\frac{1}{2}} H^{(l)} W^{(l)} \right) \quad (6)$$

Where A represents the adjacency matrix of the English knowledge graph, I_N represents the corresponding diagonal array of the English knowledge graph, and $A + I_N$ represents the self-connection added to the adjacency matrix of the English knowledge graph, i.e., assuming that there are edges connecting the exercise nodes themselves to themselves in the English knowledge graph, the English knowledge graph after adding the self-connection can avoid the loss of features in the training of the graph convolution neural network, so as to make the model better in terms of learning effect. $\check{D}^{-1/2}$ and $W^{(l)}$ denote the weight matrix to be learned in layer 1, $H^{(l)}$ denotes the state matrix of the hidden layer in layer 1, and σ denotes the activation function of the current hidden layer, which can be any activation function, and here the RELU function is chosen, and the definition of the RELU function is shown in Equation (7):

$$f(x) = \begin{cases} 0, & x \leq 0 \\ x, & x > 0 \end{cases} \quad (7)$$

The convolution of graph convolutional neural networks can be broadly categorized into two types: spectral convolution and spatial domain convolution. Spectral convolution refers to the processing of the filter and graph signals in the graph convolution network in the Fourier domain, and spatial domain convolution refers to connecting the nodes in the knowledge graph in the spatial domain, so that the nodes in the spatial domain reach a hierarchical relationship, and then later the convolution operation is carried out in the spatial domain. Equation (7) is the most basic convolution propagation in graph convolution neural network, this propagation can be replaced by using the first order approximation of the local spectral filter on the knowledge graph thus simplifying the computation, reducing the amount of computation, and speeding up the model computation speed, so in this paper we use the spectral convolution, which transfers the filter and the graph signals in the graph convolution network to the Fourier for computation.

In order to make the model have better results, in the process of spectral convolution, the feature matrix is converted to Laplace matrix L . Laplace matrix L is calculated as: $L = D - A$, where D is a diagonal matrix, each value on the diagonal is the degree of a vertex on the graph, and A refers to the adjacency matrix of the graph, which is the basic definition of Laplace matrix, and the Laplace transform of graph convolution neural network is: $L = D^{-1/2}(D - A)D^{-1/2}$. The Laplacian matrix L must be a semi-positive definite symmetric matrix with n linearly independent eigenvectors, such that the matrix is perfectly amenable to eigendecomposition, thus allowing the map convolution operation to proceed smoothly.

The spectral convolution of the graph can be defined as the product of the graph signal $x \in R^N$ and the filter $g_\theta = \text{diag}(\theta)$, which is parameterized in the Fourier domain $\theta \in R^N$ in order to facilitate the computation in the Fourier domain, as in Eq. (8):

$$g_\theta^* x = U g(\Lambda) U^T x \quad (8)$$

where $U^T x$ is the graph Fourier transform of x , $g(\Lambda)$ is a function of the eigenvalues of the graph Laplace matrix, U is a matrix consisting of the eigenvectors of the normalized graph Laplace matrix, and Λ is a diagonal array consisting of the eigenvalues of L . The graph Laplace matrix is computed as in Equation (9):

$$L = I^N \cdot D^{-1/2} A D^{-1/2} = U \Lambda U^T \quad (9)$$

The above equation can be computed on a small-scale graph, the English knowledge graph in this paper contains 9449 exercise nodes, the graph size is relatively large, the complexity of computing the eigenvalue decomposition of the graph Laplace matrix will be high, and the model computation time will become longer, in order to solve the problem, this paper uses the expansion of the k nd of the Chebyshev polynomials $T_k(x)$ to approximate instead of $g(\Lambda)$, t as in Equation (10):

$$g^\sim(\Lambda) \approx \sum_0^k \theta_k^- T_k(\Lambda^\sim) \quad (10)$$

where $\Lambda^\sim = 2\Lambda\lambda_{\max} - I_N$, $\lambda_{\max}$ denote the largest eigenvalue of L and $0^\infty \in R^k$ is the Chebyshev coefficient vector. According to the definition of Chebyshev polynomials: $T_k(x) = 2xT_{k-1}(x) - T_{k-2}(x)$, where $T_0(x) = 1, T_1(x) = x$, hence Eq:

$$g_\theta^* x \approx \sum_{k=0}^k \theta_k^- T_k \left(\frac{2}{\lambda_{\max}} L - I_N \right) x \quad (11)$$

The k st order Chebyshev polynomials used in the above equation and the Laplace polynomials in equation (8) are also of k nd order, which means that the representation of a particular node in the

English Knowledge Graph depends only on the exercise nodes with path lengths up to k . Multiple convolutional layers can be stacked in the graph convolutional neural network, and each convolutional layer consists of the above equation (12), with the increase of the number of convolutional layers, the parameters in the model increase, which leads to the emergence of the problem of slow model computation and overfitting, in order to prevent the emergence of the above problems, in this paper, the number of convolutional layers is set to 2, and its forward propagation model can be simply written as equation (12):

$$Z = f(X, A) = \text{softmax}(\tilde{A} \text{RELU}(\tilde{A}XW^0)W^1) \quad (12)$$

During model training, cross entropy is calculated as the loss function of the model on labeled data in the training set at each iteration, and the goodness of the current effect of the model is evaluated based on the value of the loss function, and the cross entropy is calculated as in equation (13):

$$L = - \sum_{l \in y_L} \sum_{f=1}^F Y_{lf} \ln Z_{lf} \quad (13)$$

Where y_L represents the index value of labeled nodes, f represents the index of all nodes, Y_{lf} is the true label of nodes, and Z_{lf} is the prediction result. Afterwards, the gradient descent algorithm is used to back-propagate, and with the iterative training of the model, the parameters are continuously adjusted to improve the model effect, which makes the categories of nodes in the personalized knowledge graph more accurate.

2.2.3. English Knowledge Graph Construction. The construction of knowledge graphs relies on data extraction from different data sources, which is the basis for subsequent applications. For English Knowledge Graph, data come from two main sources: one is the school's own data, which usually contains electronic data purchased by the school and unstructured data stored in graphic form, in which students' test scores are stored in the form of structured tables. The other is publicly available web data of foreign language articles, usually unstructured data stored in the form of web pages.

The former usually requires only simple preprocessing as input for subsequent personalization paths, but the latter usually requires the use of techniques such as natural language processing to extract unstructured information. School data provides structured information such as learners' knowledge levels and learning records, which are used to construct learner profiles and provide the basis for personalized learning and recommendation. The unstructured web corpus is rich in knowledge, and technical tools are to be used to extract entities, relationships and attributes to construct a knowledge graph.

The information provided by the two data sources helps to construct learner profiles and knowledge graphs, and lays the foundation for personalized paths for university foreign language learning based on knowledge graphs. The personalized path collects the learner's knowledge status by analyzing the learner's data and recommends personalized learning paths and content for the learner in combination with the knowledge graph. This requires preprocessing, cleaning, extracting, fusing and labeling the data from the two data sources to finally build a knowledge graph and learner user profiles.

The key to knowledge graph building lies in understanding the business and designing the knowledge graph itself. The massive resources on the Internet are an important source of information for personalized paths. As shown in Figure 1, knowledge can usually be extracted according to the syntax of subject, predicate, and object in natural language. Knowledge extraction is categorized into entity extraction, relation extraction, attribute extraction and event extraction.

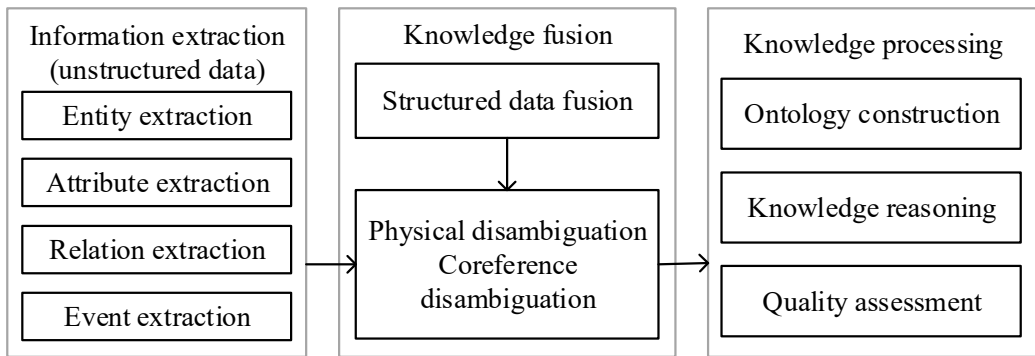


Fig. 1. Personalized path information source

2.3. Personalized Learning Path Generation

After generating the personalized English knowledge graph. According to the category of the exercise nodes, the difficulty of the exercises and the similarity between the exercises in the personalized knowledge graph, the personalized learning path is generated by drawing on the graph path generation algorithm in the data structure, and then the personalized learning path is constantly modified according to the changes in the students' learning situation, and the personalized learning path generation process is shown in Figure 2.

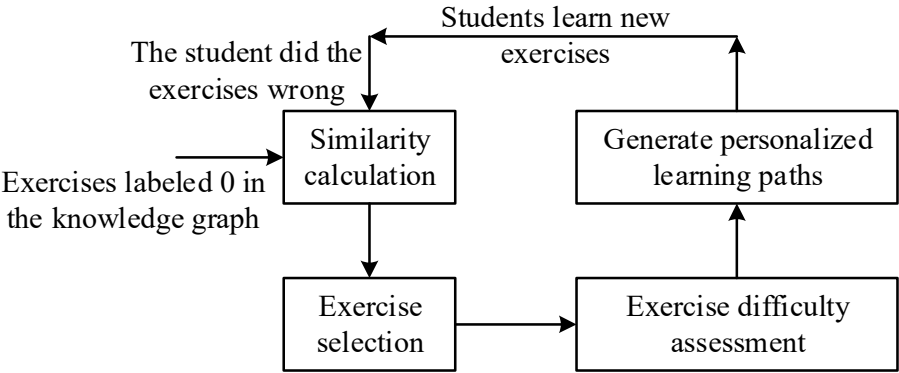


Fig. 2. Personalized learning path generation process

2.3.1. Exercise Difficulty Assessment. In order to make the generation of personalized learning paths more reasonable, it is necessary to generate a difficulty value for each exercise, and when generating personalized learning paths, the exercises are sorted according to the difficulty value of the exercises. For the assessment of the difficulty of the exercises, there is no uniform standard and norms, in this paper, when assessing the difficulty of the exercises, the difficulty of the exercises is examined from the three aspects of whether there are superlatives in the exercises, the similarity of the answers, and the number of knowledge points examined in the exercises, in which the vocabulary similarity is divided into proximity and resemblance. For essay writing, translation and other question types that cannot be assessed on the basis of the above aspects of the degree of difficulty, for these types of questions, the value of the degree of difficulty is set as the average of the degree of difficulty value of all exercises.

2.3.2. **English Learning Path Generation Based on Personalized Knowledge Mapping.** After generating the personalized knowledge graph, all the exercise nodes in the knowledge graph have labels, and the exercises that make up the personalized learning path are selected in the personalized knowledge graph by the students' doing and the similarity between the exercise nodes. Afterwards, personalized learning paths are generated based on the difficulty level of the exercises [36].

After the students have finished learning the exercises in the personalized learning path, according to the students' current doing situation, the exercises that the students have done in the English knowledge graph are marked with new markers and put into the graph convolutional neural network for training to get the newest personalized knowledge graph, and then new personalized learning paths are generated according to the generation method of personalized learning paths in the above mentioned.

2.4. Experimental analysis

2.4.1. **Experimental setup.** A total of 4,580 English grammar MCQs (single-choice questions) were used as the experimental data set. These questions were mainly from the Singapore Primary School Leaving Examination (PSLE), a selective examination administered by the Ministry of Education (MOE) for elementary school to enter secondary schools, and the Gaokao, an entrance examination for higher education schools in China, which is taken by high school graduates or those with a high school diploma. The Gaokao is a selection test for admission to higher education schools in China, taken by high school graduates or candidates with equivalent qualifications. In these experimental datasets, 400 syntactic MCQs are randomly extracted and generated as query MCQs, and the remaining 4180 MCQs will be used as questions to be retrieved from the database.

2.4.2. English Exercise Generation:

1) **Word Vector Setting.** Firstly, we need to set up the word vector pre-training method for the input layer in the GNN model. In this paper, based on the large-scale English learning corpus, the initialized word embedding vectors for each word are pre-trained using the Word2Vec method introduced in this paper, and the dimension of the word embedding vectors is 200. In particular, for the words that do not appear in the pre-training corpus of the reading comprehension text, the values of their vectors are randomly initialized.

2) **GNN model parameter settings.** In the GNN model, the maximum number of sentences M in each document was set to 25, and the maximum number of words N in each sentence was set to 40 (using 0 to fill in insufficient utterances or words). The values are set as shown in the statistical analysis in Fig. 3, i.e., 95% of the documents (utterances) contain less than 30 (50) utterances (words). For the convolutional layer setup in GNN, the model operates with 3 layers of wide convolution and 1 layer of narrow convolution. Each layer has a hidden vector dimension of (200, 400, 600, 600). The model sets the size of the convolution kernel of the convolutional layers $k = 3$ and sets the size of the pooling layer window for each layer $p = (3, 3, 2, 1)$.

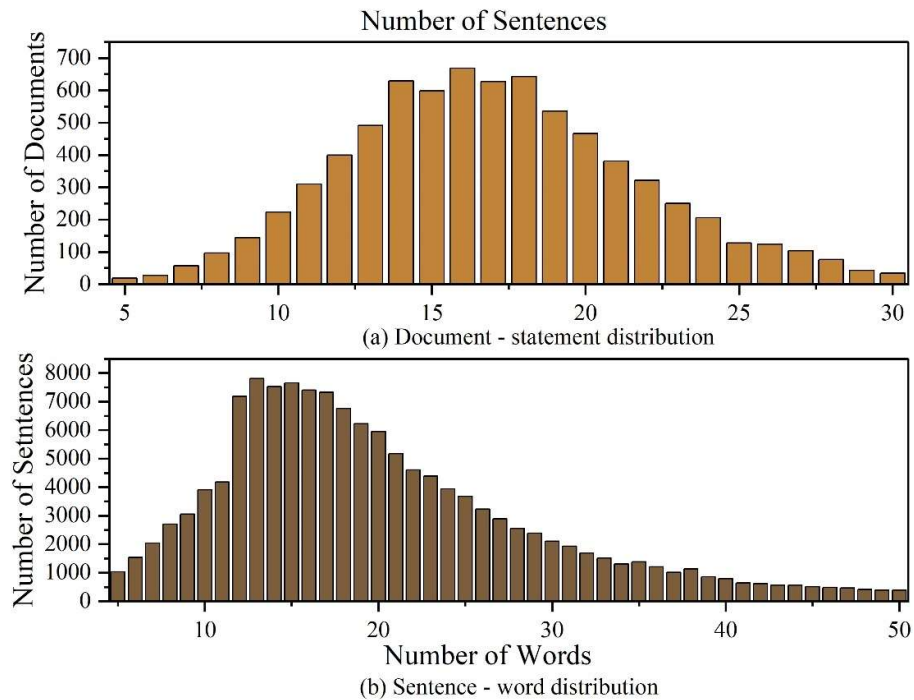


Fig. 3. Data set distribution statistics

The dataset used in this section is derived from daily practice data of English short answer questions from high school students. The short answer questions contain corresponding knowledge concepts (e.g., “grammar”). The experiment selects five of the most frequent knowledge concepts. Figure 4 shows the analysis of the short-answer question dataset. It can be seen that when the question is shorter, the proportion of formulas is bigger and when the question is only about 40 characters, the proportion of formulas is more than 3%. This phenomenon can be seen: for English short answer questions, their formulas have important information, therefore, for such questions, it is necessary to understand and represent their formulas in depth in order to effectively understand the meaning of simple questions.

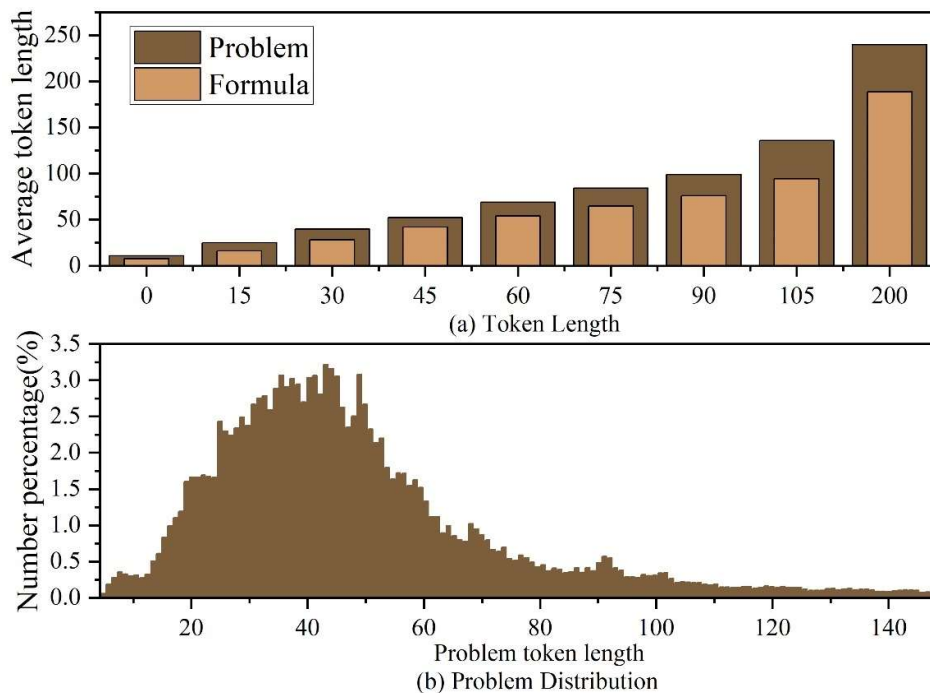
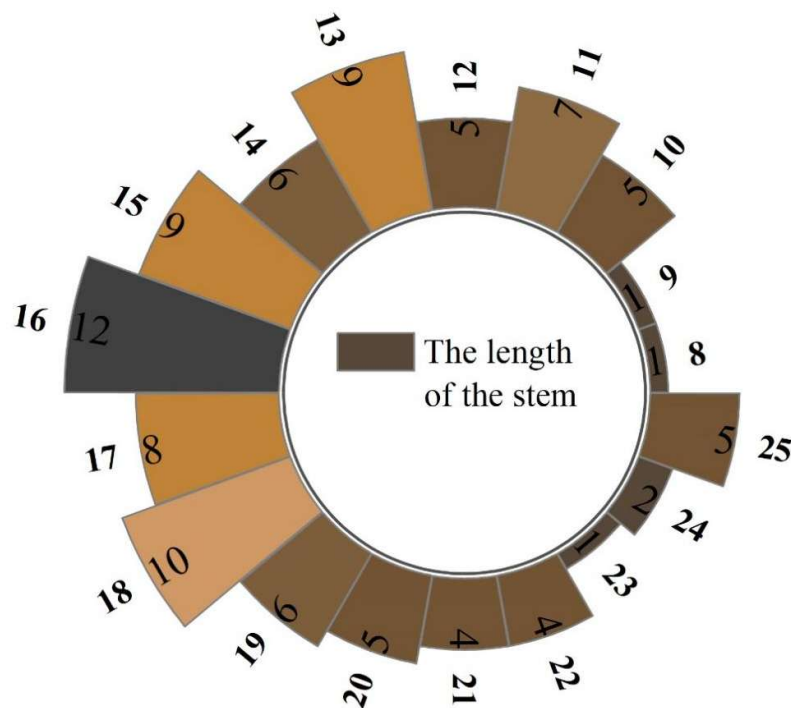


Fig. 4. Short answer data set analysis

2.4.3. Difficulty coefficient of grammatical problems. Figure 5 shows the length distribution of grammatical MCQ stems in the experimental dataset, where the average length of grammatical MCQ stems is 16.5 words. The similarity of each pair of grammatical MCQs was manually marked by 10 native English speakers, and the similarity was divided into 5 grades ("Very poor", "Fair", "Good", "Excellent", "Perfect"). Here "very poor" means that the two syntax MCQs are not related at all, and "perfect" means that the two syntax MCQs are the most correlated. In the experiment, the two syntax MCQs with "very poor" correlation were not correlated, and the other correlations were correlated. For each query MCQ, an average of 13% of the syntax MCQs in the database are related to it. 200 query MCQs were used for training and 200 query MCQs were used for testing. For each query MCQ used for training, 5 questions marked with different relevance levels were extracted from the database to learn the regression model parameters.

**Fig. 5.** The length of the syntax MCQ is distributed

English grammar questions are categorized in different ways, and each grammar type represents a grammar concept. The Grammar Concepts Statistics and Learning module counts the conceptual analysis, common methods, and grammar topics belonging to each grammar concept. In order to distinguish the difficulty level of grammar topics, this module marks each grammar topic with a difficulty label. Specifically, the Statistics and Learning of Grammar Concepts module consists of the following 3 steps:

- 1) Grammar Concept Type Statistics. Grammar Concept Type Statistics provides detailed explanations of various grammar concepts so that users can understand the meaning of each grammar concept. Commonly used grammar concepts mainly include nouns, pronouns, articles, counts, prepositions, adjectives, adverbs, conjunctions, clauses, modal verbs and verb phrases.
- 2) Labeling the Difficulty of Grammar Questions. The distribution of the difficulty coefficients of the grammar problems in the database is shown in Figure 6. Based on the difficulty coefficients, all the questions are categorized into five intervals of difficulty levels: very easy (<14), easy (14-16), average (16-18), hard (18-20) and extremely hard (>20). As can be seen in Figure 6, the difficulty

level of most of the questions is concentrated in the easy and average levels, and there are a total of 2,229 questions in these two difficulty levels, which shows that the difficulty coefficients of the English grammar questions generated by the personalized learning path are moderate.

3) Grammar Concept Learning. The user is provided with theoretical knowledge of grammatical concepts and related grammatical problems. According to the difficulty level and grammar concepts specified by the user, the eligible grammar questions are randomly selected.

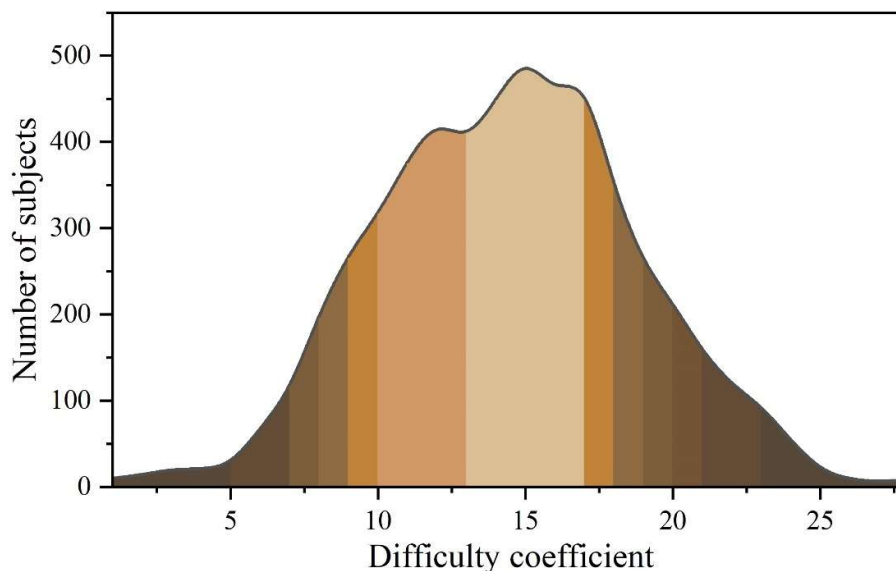


Fig. 6. Difficulty coefficient of grammar problem

3. Analysis of English Learning Effect Based on Personalized Path Generation

3.1. Analysis of the Effectiveness of English Teaching and Learning in Y Secondary School English Personalized Pathway

3.1.1. Purpose of the experiment. In order to better examine the effect of personalized pathway recommendation on English learning, and to find out whether personalized learning pathway can effectively meet the differentiated learning needs of learners and promote the achievement of the goal of learners' "learning to learn", thus ultimately facilitating the realization of personalized learning, this study investigates the impact of personalized teaching on learners' personalized learning by taking the English subject of Y Middle School as an example. In this study, we take the English subject in Middle School Y as an example, and conduct a practical investigation of the personalized teaching model to better grasp the impact of personalized teaching on learners' personalized learning.

3.1.2. Experimental subjects and methods. In order to investigate the teaching effect of personalized learning mode, this study takes the teaching mode as the independent variable, which mainly includes two kinds of traditional teaching modes and personalized teaching modes, and tries to control the interfering variables such as teacher, teaching content, students' level, teaching time, etc., and observes the dependent variable of students' learning effect, which is the achievement of the personalized learning objectives, in order to investigate the application of personalized learning mode on students' autonomy development.

The subjects of this study are students in two classes in Middle School Y, with about 50 students in each class, and the total number of samples is 100. They were divided into an experimental class that used the English learning mode recommended by the personalized pathway and a control class that

used the traditional teaching mode for English learning.

The arithmetic mean of the students' English midterm test scores in each class was used as the reference quantity in the experiment, and the overall level of the students in each class was similar, and the distribution of the students' levels within the classes was also similar. In addition, the teaching experiences and styles of these appointed teachers are similar, and the teaching contents and progress are the same, which can effectively exclude the interference of other irrelevant variables.

3.2. Learning assessment analysis

In this section, the design of evaluation indexes for personalized paths is carried out from three dimensions: satisfaction, learning attitude and learning outcome of English vocabulary learning recommendation. A total of 58 questionnaires were distributed and 50 were effectively returned.

Figure 7 shows the evaluation of the personalized path recommendation, the columns in satisfaction from left to right are using behavioral data, comprehensive resources, reasonable vocabulary recommendation, learning attitude are better tools to assist learning, improve learning interest, improve learning motivation, and the learning results are building vocabulary structure and improving learning efficiency.

61% of learners strongly support the use of learner behavioral data, 25% support the use of behavioral data, and the majority of learners support the use of behavioral data, to varying degrees, to analyze learners' personalized profiles and recommend vocabulary resources for them. 61% of learners think that the resources are very comprehensive, and 19% think that the resources are more comprehensive, and in general, the majority of learners think that the pathway recommendations in the On the whole, most learners think that the vocabulary resources in the path recommendations are comprehensive, and a few learners think that the vocabulary resources can be supplemented. At the same time, 55% of the learners and 28% of the learners think that the vocabulary recommendation meets the demand in different degrees, indicating that most learners think that the vocabulary recommendation function of the personalized path can basically meet the learners' personalized vocabulary demand, which means that the personalized path recommendation function has a good recommendation effect. Overall, most learners are satisfied with the personalized path to varying degrees, and from the analysis of the learner satisfaction survey, we know that the personalized recommendation path for English vocabulary learning has a certain degree of effectiveness.

58% of the learners and 25% of the learners think that the personalized path is a better tool to assist vocabulary learning to varying degrees, while 11% of the learners think that the personalized path is not very useful for vocabulary learning. Overall the personalized path can help the learners learn vocabulary to a certain extent, which indicates that the personalized path develops a personalized vocabulary learning sequence for the learners, and has a positive effect on the learners' learning of English vocabulary. 62% of the learners strongly agreed that personalized paths can increase learners' interest in learning English vocabulary, 21% of the learners thought that personalized paths can increase their interest in learning vocabulary, and a few learners thought that personalized paths cannot help learners to increase their interest in learning. 61% of the learners strongly agree that personalized paths can increase motivation to learn vocabulary, 21% of the learners think that personalized paths can increase motivation to learn vocabulary, a few learners think that personalized paths increase motivation to learn vocabulary to a lesser extent, and in general personalized paths have a good positive effect on learning English vocabulary.

The results of the survey show that 55% of the learners and 29% of the learners think that the personalized path can help learners to establish the overall structure of the vocabulary, which is helpful for memorizing vocabulary to a certain extent. 56% of the learners agree that the personalized path can improve the efficiency of learning English vocabulary, 22% of the learners

think that the personalized path can basically improve the efficiency of learning English vocabulary, and a few learners think that the personalized path has little or no effect on learning English vocabulary. A few learners think that the personalized path has little or no effect on English vocabulary learning. Overall, learners agree that the personalized path has some help in memorizing vocabulary.

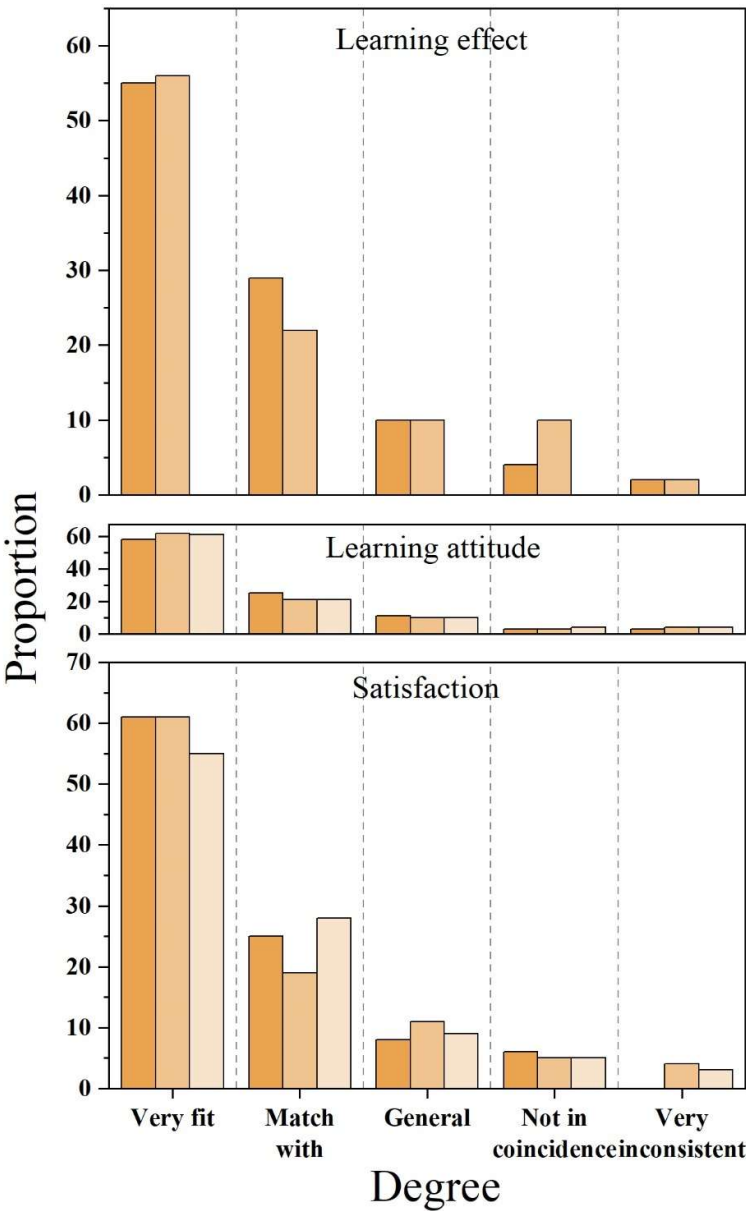


Fig. 7. Evaluation of personalized path recommendations

3.3. Comparison of Personalized Learning Effectiveness

3.3.1. Differences in learning outcomes. In this study, the experimental method was used to collect data using questionnaires on the teaching and learning of the experimental group that conducted personalized pathway teaching and the control group that taught traditionally, and SPSS 22.0 was used to analyze the data.

Table 1 shows the analysis of differences in personalized learning effects, using t-tests to test the differences in the dimensions of personalized learning goals between two different instructional models, personalized pathway instruction and traditional instruction. The results show that the

different modes of teaching and learning activities show a significant difference at 0.01 level ($p=0.00<0.01$) in the dimensions of learners' "knowledge and skills", "process and method" and "affective attitudes". ($p=0.00<0.01$), and by comparing the mean values of $3.816>3.526$, $3.815>3.096$, $3.785>3.199$, it is found that the mean value of personalized pathway teaching classes is larger than that of traditional teaching classes, which shows that compared with the traditional teaching mode, the personalized pathway teaching mode is more capable of promoting students' personalized learning. This is due to the personalized teaching mode from the problem situation, guiding students to carry out independent, cooperative inquiry learning, and focus on the independent development of students in the learning process, can stimulate the positive emotional participation of students, and better contribute to the achievement of teaching goals.

Table 1. Individualized learning effect difference analysis

Dimension	Class	Sample size	Mean value	Standard deviation	T	P
Knowledge and skills	Experimental class	50	3.816	0.759	0.049	0.000
	Comparison class	50	3.526	0.615		
Process and method	Experimental class	50	3.815	0.798	0.036	0.000
	Comparison class	50	3.096	0.604		
Emotional attitude	Experimental class	50	3.785	0.685	0.037	0.000
	Comparison class	50	3.199	0.618		

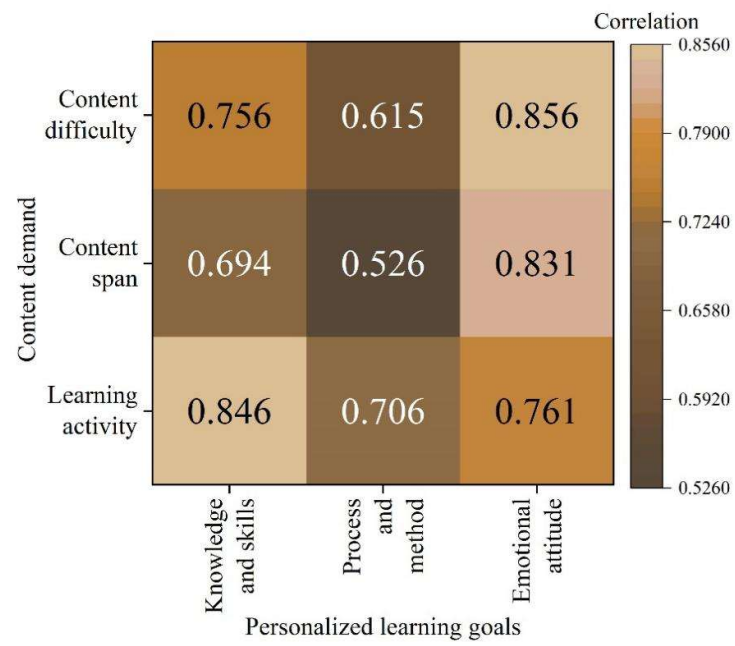
3.3.2. Relevance. The learners' content needs, resource needs, process needs, evaluation needs and personalized learning objectives are tested by correlation analysis, and the Pearson correlation coefficient is used to indicate the strength of the correlation. Figure 8 shows the results of the correlation analysis, and Figures (a)-(d) show the content needs, resource needs, process needs, and evaluation needs, respectively.

First of all, the specific correlation coefficient values of each dimension of content needs and each dimension of personalized learning objectives are all in the range of 0.5 to 0.9, indicating that there is a close correlation between learning activities and personalized learning objectives.

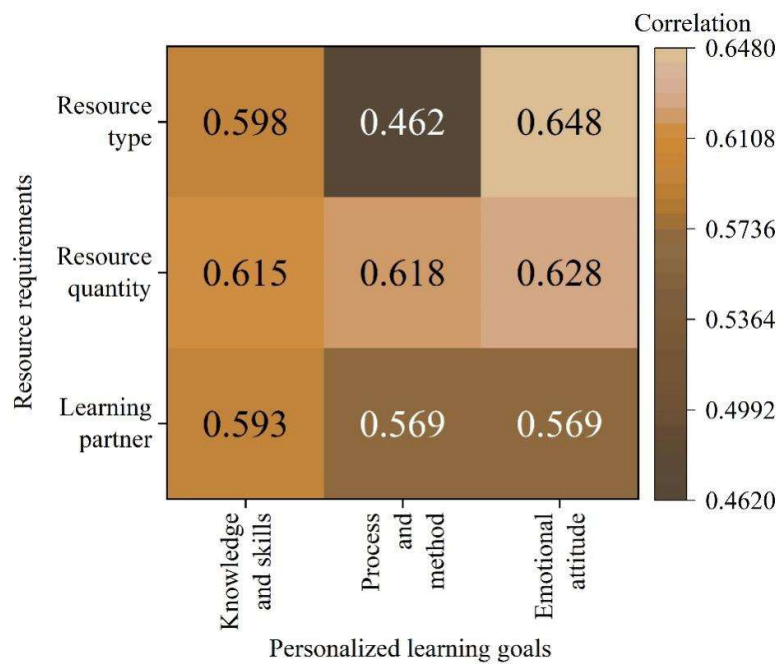
The correlation coefficients between the dimensions of resource type, resource amount and learning partner and the dimensions of personalized learning objectives are in the range of (0.4,0.7), indicating that the resource demand, including resource type, resource amount and learning partner, has a close influence on the personalized learning objectives.

The correlation coefficients between the dimensions of learning mode, learning progress, interaction mode and interaction frequency and the dimensions of personalized learning objectives are all greater than 0.6, indicating that the process needs have a close influence on personalized learning objectives, especially the high correlation with learners' emotional attitudes.

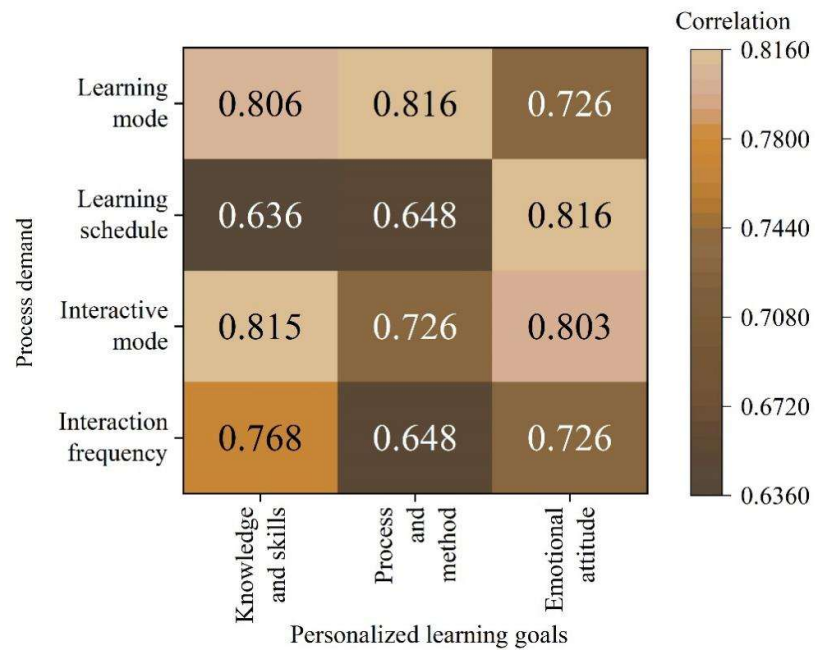
The correlation coefficient between learners' assessment needs and personalized learning objectives is $0.5<R\text{ value}<0.9$, indicating that there is a correlation between assessment needs and personalized learning objectives. The R-value of assessment methods and evaluation feedback is >0.7 , indicating that assessment methods and evaluation feedback have a close influence on the achievement of personalized learning objectives.



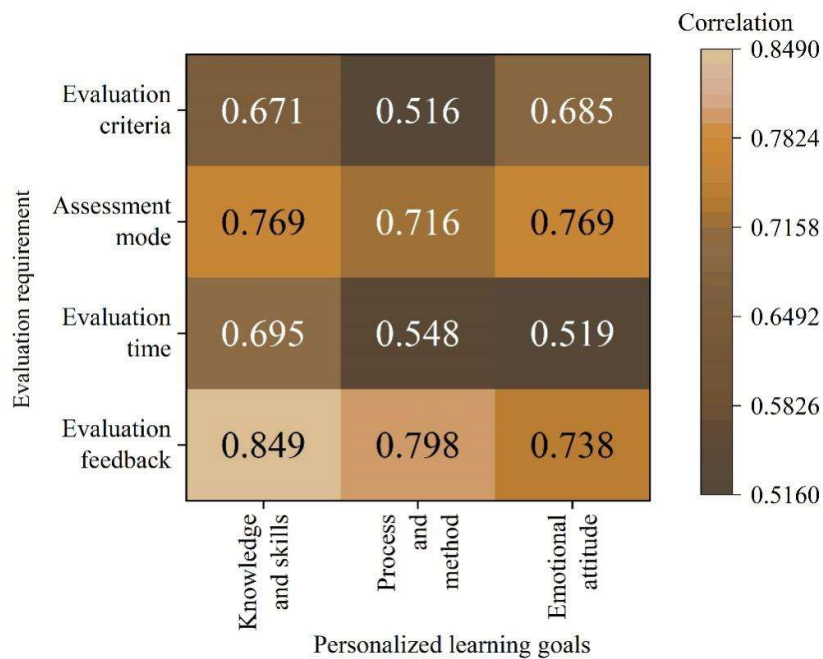
(a) Content requirements



(b) Resource requirements



(c) Process requirement



(d) Evaluation demand

Fig. 8. Correlation analysis results

3.4. Analysis of Learning Achievement

3.4.1. Level of knowledge. Changes in students' performance can directly reflect the effect of teaching, in order to exclude the influence of the differences in students' a priori knowledge level on the changes in the performance of the two groups of students, the study conducted a pre-test of the a priori knowledge level of the two groups of students, and the results of the quiz are shown in Table 2, the mean value of the experimental group is 64.299, and the mean value of the control group is 63.926, and the Sig value is 0.745 is greater than 0.05, so that the two groups of students have

comparable a priori knowledge levels are comparable, there is no significant difference and the experiment can be continued. In weeks 2-7, six English teaching activities based on personalized path generation were implemented for the students in the experimental group, and after each completion of the teaching activities of the two groups of students, the teacher carried out achievement post-tests and cognitive load tests to understand the effectiveness of the model in actual teaching.

In this study, the application of English teaching model based on personalized path generation was carried out in weeks 2-7, and six teaching activities were carried out in total, and after students completed remedial learning each time, a posttest would be conducted for the two groups of students, and the average score of the first posttest of the experimental group in the specific posttest was slightly higher than that of the control group and the Sig value was 0.346, which was greater than 0.05, which indicated that there was no significant difference in the score of the first posttest of the students in the two groups and the occurrence of this. The reason for this situation may be that the teaching content of the first lesson is less and more basic. As can be seen from the table, in the five subsequent posttests, the posttest scores of the students in the experimental group were greater than those of the control group, and in addition, the Sig values were all less than 0.05, indicating that the difference in the scores of the two groups of students was significant, and with the increase in the number of posttests, the difference between the posttest scores of the two groups of students was getting bigger and bigger, which side by side reflects the accuracy of the personalized path-generated English teaching, and the remedial learning resources pushed out to meet the needs of the students, targeted to the students' knowledge status, so the posttest scores of students in the experimental group improved more. The data show that the personalized teaching model can effectively improve students' English scores and teaching effect in the actual teaching application.

Table 2. Analysis of post-test results of experimental group and control group

/	N	Me an	Standard deviation	Minimum value	Maximum value	Sig. (Double tail)
Control group	50	64. 299	6.526	40	89	0.745
Experimental group	50	63. 926	6.218	46	85	
Serial number	Group	N	Mean	Standard deviation	Standard error	Sig. (Double tail)
Post-test score 1	Control group	50	74.299	2.168	0.348	0.346
	Experimental group	50	74.648	2.135	0.365	
Post-test score 2	Control group	50	76.469	1.816	0.309	0.046
	Experimental group	50	80.169	1.715	0.281	
Post-test score 3	Control group	50	78.961	2.036	0.366	0.013
	Experimental group	50	83.615	2.106	0.304	
Post-test score 4	Control group	50	79.615	1.496	0.249	0.004
	Experimental group	50	85.866	1.526	0.266	
Post-test score 5	Control group	50	82.816	1.278	0.159	0.000
	Experimental group	50	88.654	1.136	0.163	
Post-test score 6	Control group	50	84.826	1.146	0.149	0.000
	Experimental group	50	94.387	1.266	0.156	

3.4.2. Cognitive load. In this study, the cognitive load measurement data of the two groups of students were analyzed using SPSS 22.0, and the software generated two tables, which are the table of group statistics and the table of independent samples test, and the study summarized and simplified the two tables to obtain the results of cognitive load of the two groups of students, which are shown in Table 3.

As can be seen from the table, the mean value of the intrinsic cognitive load of the students in the experimental group is 18.936, with a standard deviation of 4.325, and the mean value of the intrinsic cognitive load of the students in the control group is 21.399, with a standard deviation of 3.318. At the same time, $P=0.002$ is much smaller than 0.05, which means that the intrinsic cognitive load of the students in the experimental group is lower than that of the control group, and the significant difference exists, which suggests that the personalized path generated by the English teaching model in actual teaching, the teacher can reasonably arrange the teaching content according to the students' knowledge level, and the students will not think that the teaching content is difficult and will not cause the students' intrinsic cognitive load to be too high. The mean values of the extrinsic cognitive load of the students in the experimental group and the control group are 16.399 and 19.248, respectively, and the standard deviations are 4.023 and 3.615, respectively, with $p=0.026$ less than 0.05, which indicates that the extrinsic cognitive load of the experimental group is lower than that of the control group, and there is a significant difference, in the process of personalized teaching, the teacher can accurately understand the knowledge status of the students, and target to solve the students' In the process of personalized teaching, teachers can accurately know the students'

knowledge status and address the students' common problems in a targeted way, and will not recommend redundant resources that do not meet the students' knowledge level, and the students are able to carry out personalized remedial learning, while the cognitive diagnostic report of the control group is not precise enough, and the teachers don't know the students' specific mastery of the knowledge points, and the explanation of the common problems of the students may increase the students' extrinsic cognitive load, and at the same time, the remedial learning resources may have redundant resources for students, which results in The external cognitive load of the control group students is too high, therefore, in the actual teaching, the personalized teaching model can effectively reduce the external cognitive load of the students. The average value of the associated cognitive load of the students in the experimental group and the control group was 29.348 and 25.826, and the standard deviation was 5.836 and 4.267, respectively, $P=0.005$ was much less than 0.05, and the associated cognitive load of the students in the experimental group was higher than that of the control group, and there was a significant difference, indicating that the personalized teaching model could help students construct new knowledge on the basis of existing knowledge in actual teaching, and the total cognitive load of students in the experimental group was smaller than that of the control group, but $P=0.526>0.05$, no significant difference. In conclusion, in actual teaching, the personalized teaching model can effectively reduce students' intrinsic cognitive load and extrinsic cognitive load, promote students to construct new knowledge on the basis of a priori knowledge, and improve the teaching effect.

Table 3. Analysis of cognitive load in experimental group and control group

/	Experimental group	Control group	T	P
Inner cognitive load	18.936±4.325	21.399±3.318	-3.154	0.002
External cognitive load	16.399±4.023	19.248±3.615	-3.439	0.026
Associated cognitive load	29.348±5.836	25.826±4.267	3.264	0.005
Total cognitive load	64.158±9.139	65.599±10.265	-0.526	0.526

4. Conclusion

In this paper, personalized path generation related technology is used to construct English knowledge map based on graph convolutional neural network, and then personalized English knowledge map is generated, and personalized learning paths are generated according to English learning related methods in personalized knowledge map. Set up experiments to evaluate the difficulty coefficients of English grammar problems of the personalized generation path.

1) The distribution of stem length of grammar single-choice questions in the experimental dataset is analyzed, and the average length is 16.5 words, meanwhile, an average of 13% of grammar single-choice questions in the database are related to it. Most of the difficulty levels of the generated English grammar exercises are concentrated in the easy and average levels, and there are a total of 2229 questions in these two difficulty levels, which shows that the difficulty coefficients of the English grammar questions generated by the personalized learning path are moderate.

2) The random survey method was used to obtain students' evaluations of personalized path recommendation for English learning in terms of three dimensions: satisfaction with English vocabulary learning recommendation, learning attitude, and learning outcome. 61% of learners strongly supported the use of learner behavior data, 61% of learners thought the resources were very comprehensive, most learners thought the vocabulary resources in the pathway recommendations were comprehensive, and a few learners thought the vocabulary resources could be supplemented.

3) The experimental method was used to analyze the difference in teaching effect between the

experimental group and the control group of traditional teaching, and the teaching activities of different modes showed significant differences of 0.01 level in the dimensions of "knowledge and skills", "process and method" and "emotional attitude", and the mean values of the experimental class were 3.816, 3.815 and 3.785, respectively, which were greater than those of the traditional teaching class.

References

- [1] Eisenstein, J. (2019). Introduction to natural language processing. The MIT Press.
- [2] Fanni, S. C., Febi, M., Aghakhanyan, G., & Neri, E. (2023). Natural language processing. In Introduction to Artificial Intelligence (pp. 87-99). Cham: Springer International Publishing.
- [3] Deng, L. (2018). Deep learning in natural language processing. Springer.
- [4] Eisenstein, J. (2018). Natural language processing. Jacob Eisenstein, 507.
- [5] Guetterman, T. C., Chang, T., DeJonckheere, M., Basu, T., Scruggs, E., & Vydiswaran, V. V. (2018). Augmenting qualitative text analysis with natural language processing: methodological study. *Journal of medical Internet research*, 20(6), e231.
- [6] Srinivasa-Desikan, B. (2018). Natural Language Processing and Computational Linguistics: A practical guide to text analysis with Python, Gensim, spaCy, and Keras. Packt Publishing Ltd.
- [7] Sarkar, D. (2019). Text analytics with Python: a practitioner's guide to natural language processing (pp. 1-674). Bangalore: Apress.
- [8] Pandey, S., & Pandey, S. K. (2019). Applying natural language processing capabilities in computerized textual analysis to measure organizational culture. *Organizational Research Methods*, 22(3), 765-797.
- [9] Ahammad, S. H., Kalangi, R. R., Nagendram, S., Inthiyaz, S., Priya, P. P., Faragallah, O. S., ... & Rashed, A. N. Z. (2024). Improved neural machine translation using Natural Language Processing (NLP). *Multimedia Tools and Applications*, 83(13), 39335-39348.
- [10] Ciampelli, S., Voppel, A. E., De Boer, J. N., Koops, S., & Sommer, I. E. C. (2023). Combining automatic speech recognition with semantic natural language processing in schizophrenia. *Psychiatry research*, 325, 115252.
- [11] Adam, E. E. B. (2020). Deep learning based NLP techniques in text to speech synthesis for communication recognition. *Journal of Soft Computing Paradigm (JSCP)*, 2(04), 209-215.
- [12] Jagannath, K. V., & Banerji, P. (2023). Personalized Learning Path (PLP)–“App” for improving academic performance and prevention of dropouts in India. *Human Interaction & Emerging Technologies (IHET-AI 2023): Artificial Intelligence & Future Applications*, 108-117.
- [13] Chen, X., Zou, D., Cheng, G., & Xie, H. (2021, July). Artificial intelligence-assisted personalized language learning: systematic review and co-citation analysis. In 2021 international conference on advanced learning technologies (ICALT) (pp. 241-245). IEEE.
- [14] Liu, Y., & Zu, Y. (2024). Design and Implementation of Adaptive English Learning System Integrating Language Contexts. *Journal of Electrical Systems*, 20(9s), 85-91.
- [15] Hernandez, J. (2024). ENHANCING LANGUAGE CIRCLES IN EARLY CHILDHOOD EDUCATION THROUGH AI INTEGRATION IN THE USA. In EDULEARN24 Proceedings (pp. 4793-4801). IATED.
- [16] Xu, G., & Wong, C. U. I. (2024). Deep Learning-Based Educational Image Content Understanding and Personalized Learning Path Recommendation. *Traitement du Signal*, 41(1).
- [17] Kassimi, M. A., & Essayad, A. (2022, May). Natural Language Processing and Motivation for Language

- Learning. In International Conference on Advanced Intelligent Systems for Sustainable Development (pp. 295-307). Cham: Springer Nature Switzerland.
- [18] Tapalova, O., & Zhiyenbayeva, N. (2022). Artificial intelligence in education: AIED for personalised learning pathways. *Electronic Journal of e-Learning*, 20(5), 639-653.
 - [19] Agnnes, & Jureynolds. (2024, August). The Effectiveness of Personalized Learning E-Platform to Support Sustainable Education. In 2024 3rd International Conference on Creative Communication and Innovative Technology (ICCIIT) (pp. 1-7). IEEE.
 - [20] Xia, Y., Shin, S. Y., & Shin, K. S. (2024). Designing Personalized Learning Paths for Foreign Language Acquisition Using Big Data: Theoretical and Empirical Analysis. *Applied Sciences*, 14(20), 9506.
 - [21] Li, X., & Zhang, B. (2024). Personalized Learning Path Recommendation Algorithm for English Listening Learning. *Journal of Electrical Systems*, 20(6s), 2188-2199.
 - [22] Wu, J. (2024, April). Design and Evaluation of Personalized English Learning System Based on Artificial Intelligence. In Proceedings of the International Conference on Decision Science & Management (pp. 276-281).
 - [23] Pan, Y. (2024, May). Personalized English Vocabulary Learning Path Recommendation Based on Reinforcement Learning. In Proceedings of the 2024 International Conference on Machine Intelligence and Digital Applications (pp. 529-533).
 - [24] Lu, X., & Samah, N. (2024). Research on Intelligent Educational Technology in Business English Education in Building Multi-model Learning Environments and Personalized Learning Paths. *Journal of Human Centered Technology*, 3(2), 84-92.
 - [25] Arani, S. M. N. (2024). Navigating the future of language learning: A Conceptual review of ai's role in personalized learning. *Computer-Assisted Language Learning Electronic Journal*, 25(3), 1-22.
 - [26] Shakhnoza, O., & Nargiza, I. (2024). NATURAL LANGUAGE PROCESSING (NLP) APPLICATIONS IN LANGUAGE TEACHING. «CONTEMPORARY TECHNOLOGIES OF COMPUTATIONAL LINGUISTICS», 2(22.04), 278-283.
 - [27] Zhang, M. (2024). Enhancing MALL with Artificial Intelligence: Personalized Learning Paths in EFL Teaching. *Research and Advances in Education*, 3(8), 49-61.
 - [28] Yenuri, A. (2024). AI-Powered Language Learning: A New Frontier in Personalized Education. *Journal of Language Instruction and Applied Linguistics*, 1(02), 38-45.
 - [29] Sangkala, I., & Mardonovna, N. S. (2024). ARTIFICIAL INTELLIGENCE AS A PERSONALIZED TUTOR IN LANGUAGE LEARNING: A SYSTEMATIC REVIEW. *KLASIKAL: JOURNAL OF EDUCATION, LANGUAGE TEACHING AND SCIENCE*, 6(2), 565-576.
 - [30] Zhou, Z., & Li, W. (2024). Personalized College English Learning Based on Artificial Intelligence: Algorithm-Driven Adaptive Learning Method. *International Journal of High Speed Electronics and Systems*, 2540102.
 - [31] Ismayilli, F. (2024). REVOLUTIONIZING LANGUAGE LEARNING: INTEGRATION OF AI TO PERSONALIZED LANGUAGE LEARNING. *German International Journal of Modern Science/Deutsche Internationale Zeitschrift für Zeitgenössische Wissenschaft*, (79).
 - [32] Yue Li & Endong Wang. (2024). Construction of personalized recommendation model for educational video game resources based on knowledge graph. *Entertainment Computing* 100660-.
 - [33] Gu JiaKai, Li, Vo D. Nam & Jung J. Jason. (2022). Contextual Word2Vec Model for Understanding Chinese Out of Vocabularies on Online Social Media. *International Journal on Semantic Web and Information Systems (IJSWIS)*(1), 1-14.

- [34] Tao Jie. (2022). Optimization of Dynamic Graphic Packaging Design Scheme Based on Graph Neural Network. *Mobile Information Systems*.
- [35] Yao Zhang. (2024). Research on the Construction of a Diversified System of Preschool Physical Education Curriculum Based on Kruskal Algorithm. *Applied Mathematics and Nonlinear Sciences*(1),
- [36] Zhaoyu Shou,Yiru Wen,Pan Chen & Huibing Zhang. (2021). Personalized Knowledge Map Recommendations based on Interactive Behavior Preferences. *International Journal of Performability Engineering*(1),36-49.