

A neural network-based model study of dynamic learning path generation for blended teaching of English majors

Yilin Wu¹, 

¹ School of Foreign Languages, Pingxiang University, Pingxiang, Jiangxi, 337000, China

ABSTRACT

Traditional learning path planning methods often fail to meet the individualized needs of learners. In this paper, a dynamic learning path planning method based on neural network is studied by constructing a student model. Firstly, the construction method of the student model is designed, and the Item Response Theory (IRT) is used as a test method for the cognitive level of students, which realizes the dynamic acquisition of student information. A neural network-based cognitive collaborative filtering model was constructed, which models learners' learning behaviors and interests, and customizes personalized dynamic learning paths for learners after assessing their cognitive levels and learning difficulties. The collaborative filtering algorithm in this paper performs better than the other four algorithms in terms of accuracy and coverage, and the accuracy and coverage rate of the generated knowledge point sequences reach 98.9% and 93.6% respectively, and the performance of the students in the experimental group has been significantly improved under the application of the dynamic learning path generation model of blended teaching in this paper, indicating that the effectiveness and feasibility of the personalized learning path generation model in this paper are excellent and are expected to be further promoted.

Keywords: Neural Networks, Collaborative Filtering, Item Response Theory, Student Cognitive Level Testing, Personalized Learning Pathways

1. Introduction

In the brand-new era of rapid development of information technology, artificial intelligence technology has become an important pusher of educational change, showing great potential in the teaching and learning process of college English [1]. Especially with the breakthroughs in the field of artificial intelligence such as speech recognition and natural language processing, university English teaching is gradually expanding the use of these advanced technologies, effectively promoting the renewal of educational thinking and the change of teaching mode [2]. Hybrid teaching mode is a teaching method

 Corresponding author.

E-mail address: Maggie241006@126.com (Y. Wu).

Received 05 January 2024; Revised 12 February 2024; Accepted 15 November 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-514](https://doi.org/10.61091/jcmcc127b-514)

© 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/>).

combining traditional classroom teaching with modern network technology, which can provide students with more flexible personalized learning programs. With the assistance of artificial intelligence, this mode can not only enhance the interactivity and interest of teaching, but also expand the spatial and temporal boundaries of learning, greatly enriching the teaching resources and learning forms of university English courses [3-4].

At the same time, the hybrid teaching mode combines the advantages of online and offline teaching, and through the empowerment of intelligent technology, it can provide a more flexible and personalized teaching experience to help students effectively acquire English knowledge and skills [5]. With the development of intelligent technology, hybrid teaching mode will gradually become a new trend [6]. Therefore, it has become crucial for educators to explore in depth how to effectively integrate artificial intelligence technology and university English blended teaching, and optimize teaching resources and methods to improve the quality of teaching and learning effects [7].

Literature [8] envisioned to look at an English hybrid intelligence-assisted teaching model with a mobile information system as the underlying architecture, which can be based on student learning for targeted curriculum and teaching content design, and combined with an offline teaching model for English teaching. Literature [9] studied the current situation of the practice of mixed teaching mode of English in higher vocational education around the problems of single form of multimedia teaching and low utilization of information technology tools, and concluded that it is necessary to scientifically integrate the traditional teaching classroom and the flipped classroom, and try to provide students with an independent learning environment based on the network. Literature [10] discusses the innovative direction of English teaching in colleges and universities as how to effectively combine Internet information technology and traditional teaching classroom to promote the reform of informationization and intelligence of English teaching, focusing on the analysis of the integration of the blended teaching mode as well as practice. Literature [11] based on the Internet of Things platform, created a hybrid teaching method of college English, compared with offline English teaching methods, effectively improve the English level of students, the study provides an important reference for the development of hybrid courses in college English. Literature [12] carries out a teaching comparison experiment reveals that the English multimedia teaching mode based on Internet computer-aided technology relieves students' English learning pressure to a certain extent, and students express a definable attitude towards the new teaching mode, and believe that this scaffolding teaching mode based on Internet computer-aided technology can effectively stimulate students' enthusiasm for English learning. Literature [13] conceptualized an experimental teaching aid system with Android platform as the core logic, which can support the online and offline hybrid teaching mode of English, and based on the practical feedback, it was learned that the system promotes the development of students' thinking in the process of English learning, and cultivates students' ability of interactive communication. Literature [14] elaborated that the development of information technology promotes the innovation, upgrading and reform of higher vocational education, in which information technology-enabled higher vocational education has a positive effect on stimulating students' motivation to learn, and therefore analyzed in depth the path of online-offline hybrid English teaching mode based on information technology. Literature [15] found in the comparative analysis of hybrid English teaching classroom and traditional English teaching classroom that the effect of hybrid English teaching is better than that of traditional teaching classroom, pointing out that the hybrid English teaching mode not only promotes the provision of student's English performance, but also helps students to develop their independent learning ability and collaborative ability.

This paper introduces the IRT test model to test students' cognitive level and constructs a student test bank with reference to the construction standard of student model. Then a personalized dynamic learning path generation model based on neurocognitive collaborative filtering algorithm is proposed, which integrates students' dynamic learning behaviors and static behavioral parameters to generate personalized dynamic learning paths for students. Simulation experiments on the model

recommendation performance and experiments on the application effect of the model are conducted to demonstrate the utility of the method designed in this paper in generating personalized learning paths.

2. Dynamic learning path generation model based on neural network

2.1. Student model

2.1.1. Student modeling methodology. A student model is a data structure used to characterize the current knowledge state of a learner, reflecting the individual characteristics, knowledge state, and cognitive ability of a student. The student data in the student model is an important basis for the system to make content recommendations. The student information is categorized into eight categories: including personal information, academic information, relationship information, performance information, preference information, portfolio information, management information, and security information. The model consists of two parts: basic student information and knowledge state. The basic information is used to construct the static student model, and the knowledge state is used to carry out the dynamic history of the model through the data of student behaviors, cognitive levels, and learning paths. The specific construction process is as follows:

Step1: Basic information is collected after the first registration, which is used to build the static student model on the one hand, and to generate the credentials for logging into the system entrance on the other hand.

Step2: After logging into the system, students can choose the content to learn and test, and student behavior data is collected during the learning process, which is used to record the current learning progress of students and update the knowledge status data, so as to facilitate the students to log into the system again for the next step of learning.

Step3: When students finish learning the current content, they need to answer the corresponding test questions. The system utilizes the Item Response Theory testing model to test the cognitive level, which is an important influencing factor for updating the knowledge status data.

Step4: If students have already learned certain content, they can directly enter the test session to avoid wasting time.

Step5: After students finish the learning or testing session, the model is dynamically updated by combining the learning behavior and cognitive level data.

Step6: Finally, the updated model data is stored in the student database, which is convenient for the system to call subsequently.

2.1.2. Cognitive level testing method based on IRT theory. In previous testing sessions, students' mastery of knowledge points, i.e., cognitive level, is often measured by test scores. Item Response Theory can make effective use of data from testing sessions by studying the relationship between ability level, item difficulty, item differentiation, and guessing coefficients.

In order to ensure that the testing session can generate more valuable learning data, we need to develop more effective test items. The validity of the test questions is characterized in the IRT model by an information function, where a larger value of the function represents a more accurate estimation of the ability level. Its mathematical expression is: $I_i(\theta) = \frac{(P_i)^2}{P_i(1-P_i)}$. For the three-parameter Logistic model [16] the information function can again be expressed as:

$$I_i(\theta) = \frac{1.7^2(a_i)^2(1-c_i)}{\left[c_i + e^{1.7a_i(\theta-b_i)}\right]\left[1 + e^{-1.7a_i(\theta-b_i)}\right]^2} \quad (1)$$

where a_i, b_i, c_i, θ has the same meaning as the parameters in the Logistic model. Item response theory holds under the assumption of localization, so the test items are designed to be non-interfering and independent of each other. At this time, the test information function is the sum of all information functions, i.e.: $I(\theta) = \sum_{i=1}^m I_i(\theta)$, m is the total number of test items, which is an important parameter used to identify the overall performance of the test items.

Metrics. Estimated standard error is recorded as $SE(\theta)$, the mathematical expression is:

$$SE(\theta) = \frac{1}{\sqrt{\sum I_i(\theta)}}, \text{ and then you can get the relationship between the test information function and}$$

the estimated standard error, the greater the amount of information obtained from the test, the smaller the estimated standard error.

In summary, the following test item design rules are summarized:

Rule 1: Based on the local hypothesis conditions, all test items try to be independent of each other and unrelated to each other.

Rule 2: The larger the number of test items, the larger the value of the overall information function and the smaller the estimated standard error.

Rule 3: As can be seen from the item characteristic curve, the greater the differentiation of the items, the easier it is to differentiate between the ability levels of the subjects, and the better the quality of the test.

When the subject's ability level is slightly higher than the difficulty of the item, the more information the test item provides, the more accurate the test result is.

Cognitive level testing process

Step1: During the construction of the test bank, try to follow the design rules in the previous section to select representative test questions, then let students answer them, and calculate the difficulty coefficients and differentiation degrees of the test questions according to their answers.

Step2: Before conducting the cognitive water-half test, the initial water-half is assigned to the students, who are first allowed to answer a certain number of test questions, and then the natural logarithm of the ratio of the subjects' scores to the lost scores in the test is calculated to find out.

Step3: During the testing process, test questions that are comparable to the current cognitive level of the subjects are first presented, and the water-half value of the subjects is reassessed according to the answering situation and the termination condition is determined, and the relevant information is outputted when the termination condition is satisfied.

Step4: When the termination condition is not met, push the test question according to the answering situation. When the previous question is answered incorrectly, a slightly easier test question is pushed, otherwise a more difficult test question is pushed.

The specific realization flow is shown in Figure 1:

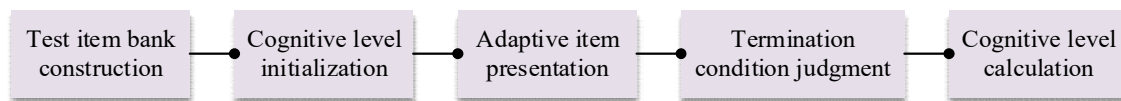


Fig. 1. Flow chart of cognitive level test

By analyzing the method of calculating the values of each ability in item response theory, the cognitive level of the tester is estimated under the assumption that the parameters of the test question are known. The method of realization is as follows:

Assuming that N student responds to m test items, the α rd student's cognitive level is noted as $\theta_\alpha (1 \leq \alpha \leq N)$, $a_i, b_i, c_i (1 \leq j \leq m)$. The differentiation, difficulty, and guessing coefficients for the j th test item, respectively, are set to 0-1 ratings for all test items and are noted as:

$$U_{\alpha j} = \begin{cases} 1 & \text{Students with ability score } \theta_{\alpha} \text{ answered question } j \text{ correctly} \\ 0 & \text{Student with an ability score of } \theta_{\alpha} \text{ answered question } j \text{ incorrectly} \end{cases} \quad (2)$$

The responses of students with a proficiency of θ_{α} on the m items are shown below:

$$U_{\alpha} = (U_{\alpha 1}, U_{\alpha 2}, \dots, U_{\alpha m}) (1 \leq \alpha \leq N) \quad (3)$$

Then the responses of N student to m test questions can be represented by matrix U :

$$U = (U_{\alpha j})_{N \times M} = \begin{bmatrix} U_{11} & U_{12} & \dots & U_{1m} \\ U_{21} & U_{22} & \dots & U_{2m} \\ U_{31} & U_{32} & \dots & U_{3m} \\ \dots & \dots & \dots & \dots \\ U_{N1} & U_{N2} & \dots & U_{Nm} \end{bmatrix} \quad (4)$$

where $U_{\alpha j}$ takes the value of 0 or 1 for a random variable whose observation is $u_{\alpha j}$. Let $P_{\alpha j}$ be the probability that a subject with ability θ_{α} answers the j th test question correctly, i.e.:

$$P_{\alpha j} = P(u_{\alpha j} = 1) \quad (5)$$

Assuming that the test questions satisfy the local independence assumption condition, the probability of various response scenarios for students with cognitive level θ_{α} respectively in a test with test question parameter $a_j, b_j, c_j (1 \leq j \leq m)$ can be expressed as L , i.e.:

$$L = \prod_{\alpha=1}^N \prod_{j=1}^m P_{\alpha j}^{u_{\alpha j}} (1 - P_{\alpha j})^{1-u_{\alpha j}} \quad (6)$$

Solve using the great likelihood estimation method. The value of the independent variable for which the likelihood function L takes a great value is taken as the value of the cognitive level to be estimated. Since L and $\ln L$ have the same point of great magnitude, taking the logarithm for L yields:

$$\ln L = \sum_{\alpha=1}^N \sum_{j=1}^m (U_{\alpha j} \ln P_{\alpha j} + (1 - U_{\alpha j}) \ln(1 - P_{\alpha j})) \quad (7)$$

Order:

$$\ln L \quad (8)$$

The partial derivatives for each parameter are equal to zero, i.e:

$$\begin{aligned} \frac{\partial \ln L}{\partial \theta_{\alpha}} &= 0 (1 \leq \alpha \leq N) \\ \frac{\partial \ln L}{\partial a_j} &= 0 (1 \leq j \leq m) \\ \frac{\partial \ln L}{\partial b_j} &= 0 (1 \leq j \leq m) \\ \frac{\partial \ln L}{\partial c_j} &= 0 (1 \leq j \leq m) \end{aligned} \quad (9)$$

Further organizing has to be done:

$$\begin{aligned}
\sum_{j=1}^m \frac{(U_{\alpha j} - P_{\alpha j})}{P_{\alpha j}(1 - P_{\alpha j})} \frac{\partial P_{\alpha j}}{\partial \theta_{\alpha}} &= 0 \\
\sum_{\alpha=1}^N \frac{(U_{\alpha j} - P_{\alpha j})}{P_{\alpha j}(1 - P_{\alpha j})} \frac{\partial P_{\alpha j}}{\partial a_j} &= 0 \\
\sum_{\alpha=1}^N \frac{(U_{\alpha j} - P_{\alpha j})}{P_{\alpha j}(1 - P_{\alpha j})} \frac{\partial P_{\alpha j}}{\partial b_j} &= 0 \\
\sum_{\alpha=1}^N \frac{(U_{\alpha j} - P_{\alpha j})}{P_{\alpha j}(1 - P_{\alpha j})} \frac{\partial P_{\alpha j}}{\partial c_j} &= 0
\end{aligned} \tag{10}$$

The above equation is a nonlinear equation about cognitive level θ_{α} . The Newton-Raphson iterative method is utilized to find the estimation of θ_{α} provided that the parameters of each item are known, denoted $g(\theta_{\alpha})$:

$$g(\theta_{\alpha}) = \frac{\partial \ln L}{\partial \theta_{\alpha}} = \frac{\partial \ln L}{\partial P_{\alpha j}} \frac{\partial P_{\alpha j}}{\partial \theta_{\alpha}} = \sum_{j=1}^m \frac{(U_{\alpha j} - P_{\alpha j})}{P_{\alpha j}(1 - P_{\alpha j})} \frac{\partial P_{\alpha j}}{\partial \theta_{\alpha}} = 0 \tag{11}$$

In the three-parameter logistic model $g(\theta_{\alpha})$ can again be expressed in the following form:

$$g(\theta_{\alpha}) = \sum_{j=1}^m D a_j \frac{1 - c_j}{P_{\alpha j} - c_j} \frac{(u_{\alpha j} - P_{\alpha j})}{P_{\alpha j}} = 0 \tag{12}$$

Newton-Raphson iteration has three elements: iteration formula, iteration initial value, and iteration termination rule. In this paper the iteration formula for the problem solved by cognitive level is:

$$\theta_{\alpha(k+1)} = \theta_{\alpha k} - g(\theta_{\alpha k}) / g'(\theta_{\alpha k}) \tag{13}$$

The expression for $g'(\theta_{\alpha k})$ is obtained by derivation:

$$g'(\theta_{\alpha k}) = \sum_{j=1}^m D^2 a_j^2 (1 - c_j)(P_{\alpha j} - c_j)(1 - P_{\alpha j}) \tag{14}$$

The initial value of the iteration is the first step in the solution of the equation, where the initial value of the cognitive level is $\theta_{\alpha 0}$ is found by the natural logarithm of the ratio of the subject's scores on the test to the missing scores, viz:

$$\theta_{\alpha 0} = \ln \frac{\sum_{j=1}^m u_{\alpha j}}{m - \sum_{j=1}^m u_{\alpha j}} \tag{15}$$

The send generation termination rule is shown in the following equation:

$$|\theta_{\alpha(k+1)} - \theta_{\alpha k}| < \varepsilon \tag{16}$$

Above that is the solution process and working principle of IRT cognitive level prediction model, Figure 2 shows the specific workflow of IRT test model.

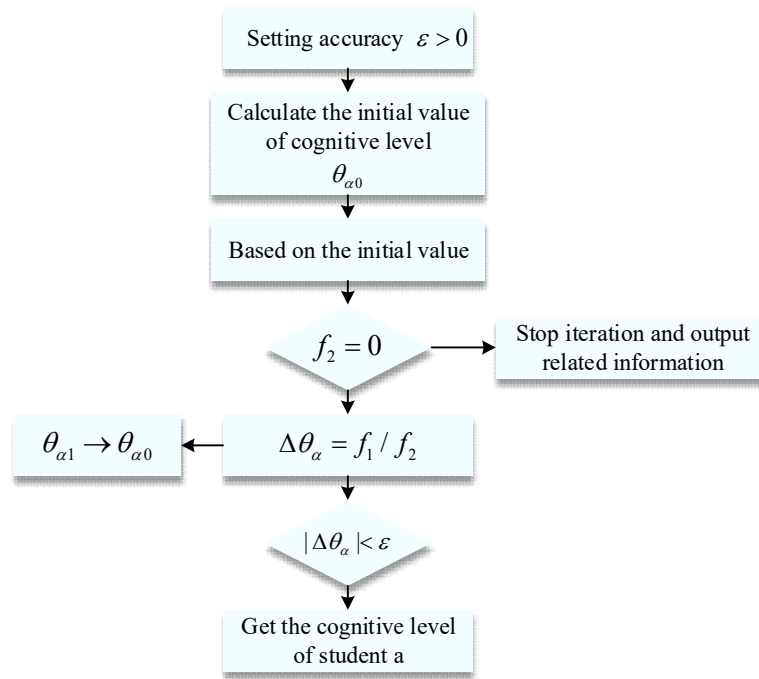


Fig. 2. Workflow of IRT cognitive level estimation model

2.2. Construction of a neurocognitive collaborative filtering model

2.2.1. Neural Collaborative Filtering Layer. The core idea of collaborative filtering [17] is to discover the correlation between learners and knowledge points based on learner behavioral data for personalized recommendation. Its core can be summarized into two main aspects, which are learner similarity or knowledge point similarity computation [18] and prediction of ratings or generation of recommendation lists.

1) Similarity Calculation. The calculation of similarity is unified in this paper by using Euclidean distance to measure the similarity relationship between the two.

In this paper, all the learner embedding vectors are weighted and summed with the knowledge point embedding vectors to calculate the fusion coding of knowledge points:

$$F_j = \sum_{i=1}^m \alpha_{ij} \cdot \sigma(U_i \cdot P_j^*) \quad (17)$$

where F_j denotes the fusion coding of knowledge point j , α_{ij} is the weight of learner i on knowledge point j , U_i is the embedding vector of learner i , P_j is the embedding vector of knowledge point j , and $\sigma(x) = \frac{1}{1 + e^{-x}}$.

In this paper, we use a similarity matrix to represent the similarity between learners and adjust the target learner's encoding by weighted summing the encodings of similar learners, where s_{ij} represents the similarity between learner i and learner j . This similarity matrix is then used to compute the new encoding of the learner vector:

$$U_{i'} = \sum_{j=1}^m S_{ij} \cdot U_j + \sum_{k=1}^n \gamma_{ik} \cdot F_k \quad (18)$$

where, U'_i denotes the new coding of learner i , S_{ij} denotes the similarity between learner i and learner j , U_j denotes the original coding of learner j , F_k is the fusion coding denoting knowledge point k , and γ_{ik} is the weight associated with learner i and knowledge point k .

2) Generation of Predictive Scoring Matrix. In the field of education, collaborative filtering-based methods can be applied to predict learners' ratings of unevaluated knowledge points in order to generate a personalized list of knowledge point recommendations. This approach bases its recommendations on the similarity between learners and the similarity between knowledge points.

Suppose, the sequence of learner feature vectors can be denoted as $U_1, U_2, \dots, U_N$ and the sequence of knowledge point feature vectors can be denoted as $P_1, P_2, \dots, P_N$. Then the rating of learner i on knowledge point j can be predicted by the similarity between the feature vector U_i of learner i and the feature vector P_j of knowledge point j . This can be represented by equation (19):

$$\hat{r}_{ij} = \left(\sum_{k=1}^N U_{ik} \cdot P_{jk} + b \right) \cdot E \quad (19)$$

where b is the bias term and E is the unit matrix.

2.2.2. Cognitive level diagnostic layer. The cognitive level diagnostic layer focuses on analyzing and decomposing the eigenvalues of learners and knowledge points to assess learners' mastery of specific aspects of different knowledge points. This part is divided into two main parts: static parameter estimation and dynamic data analysis.

1) Static parameter estimation. In the static parameter estimation part, U_{ik} denotes the learner's i response to the test k and P_{kj} denotes the test k examination of the knowledge points j so that the learner's i mastery of the knowledge points j can be expressed as a score matrix y_{ij} as in Equation (20):

$$y_{ij} = \begin{bmatrix} S_{11} & S_{12} & \cdots & S_{1j} \\ S_{21} & S_{22} & \cdots & S_{2j} \\ \vdots & \vdots & \ddots & \vdots \\ S_{i1} & S_{i2} & \cdots & S_{ij} \end{bmatrix} \quad (20)$$

In order to identify students' apparent weak knowledge points, the problem can be decomposed into different knowledge points or combinations of knowledge points, and a binary variable can be defined for each knowledge point or combination of knowledge points to indicate whether the student has mastered the knowledge point or not. Assuming that Topic k belongs to Knowledge Point j , then element S_{ij} in Score Matrix y_{ij} can be defined as a binary variable indicating whether the student i has mastered Knowledge Point j or not, as in Equation (21):

$$S_{ij} = \begin{cases} 1, & \text{if } \sum_k U_{ik} \times P_{kj} > b_{ij} \\ 0, & \text{otherwise} \end{cases} \quad (21)$$

In order to further balance and reduce the errors in the process of learners' static parameter estimation and dynamic data analysis, the static parameter estimation feature parameter b_{ij} is introduced here as the a priori information for the dynamic data analysis, which is defined as Equation (22):

$$b_{ij} = b_i + b_j = \frac{1}{J} \times \sum_{j=1}^J S_{ij} + \frac{1}{I} \times \sum_{i=1}^I S_{ij} \quad (22)$$

Therefore, it is only necessary to filter out the first q knowledge points that the proportion of knowledge points mastered by this learner is less than the proportion of knowledge points mastered by all learners, i.e., the explicitly weak knowledge point set learners of this learner, which can be expressed as Equation (23):

$$X_{ij}(i) = \left\{ j : S_{ij} = 1, \text{rank} \left(\frac{b_i}{\sum_{i=1}^I S_{ij}} \right) < \frac{1}{q} \cdot \sum_{i=1}^q S_{ij} \right\} \quad (23)$$

2) Dynamic parameter analysis. In the process of dynamic learning data analysis, the learner-knowledge point feature vector interaction matrix output from the neural collaborative filtering layer is weighted and averaged over its components sequentially to obtain the interaction matrix R_{ij} , which can be expressed as Equation (24):

$$R_{ij} = \sigma \left(w_1 \cdot \frac{1}{n} \sum_{i=1}^n (\beta_{1,i} \cdot d_{ij}) + w_2 \cdot \frac{1}{n} \sum_{i=1}^n (\beta_{2,i} \cdot c_{ij}) \right) \quad (24)$$

Next, in this paper, the interaction matrix R_{ij} is decomposed by SVD, and the learner's low-dimensional feature projection matrix R_i and the learner's behavioral data latent feature matrix R_j can be expressed as Eq. (25):

$$R_{ij} = \sum_{k=1}^K R_{ik} \cdot R_{jk} + \hat{\epsilon} \quad (25)$$

where K is the dimension of the low-dimensional feature, R_{ik} denotes the projection value of learner i on the k th low-dimensional feature, R_{jk} denotes the distribution of weight value knowledge points j on the k th low-dimensional feature, and $\hat{\epsilon}$ denotes the error matrix. Based on this, in this paper, R_i and R_j are considered as low-dimensional representations of learner and behavioral data.

As a result, the set of implicitly weak knowledge points Z_{ij} can be represented by a quintuple that satisfies Equation (26):

$$Z_{ij}(i) = \left\{ j : \max_{\theta} \sum_{(i,j) \in \square} \log p(z_{i,j} | b_{i,j}, R_i, R_j, \sigma_R, \theta_k) \right\} \quad (26)$$

where $\square$ denotes the target space, θ denotes the model parameters, $p(Z_i, j | b_{i,j}, R_i, R_j, \sigma_R, \theta)$ denotes the probability that a learner's i performance on a knowledge point j is predicted to be weak, and σ_R denotes the learner's historical learning behavior data.

2.2.3. Learning path generation layer. Based on the above analysis, the basic characteristic parameter information of this learner has been obtained through the neural collaborative filtering layer and the cognitive level diagnosis layer, therefore, the fuzzy cognitive diagnosis result of the learner under the model can be expressed as Equation (27)

$$\hat{y}_{ui} = \rho \cdot X_{ij} + (1 - \rho) \cdot Z_{ij} \quad (27)$$

where $\rho \in [0,1]$ represents the ratio parameter between the common features and the learner-independent attribute mastery patterns.

Since X_{ij} is a static covariate that remains virtually unchanged during the learning session interaction recordings, only the dynamic parameter estimation component, Z_{ij} , needs to be updated and self-healed. Its final objective function can be expressed as equation (28):

$$L = \frac{1}{2} \sum_{(i,j) \in R_{ij}} (R_{ij} - \beta_1^T \beta_2)^2 + \frac{\lambda_1}{2} \|\beta_1\|^2 + \frac{\lambda_2}{2} \|\beta_2\|^2 \quad (28)$$

where λ_1 and λ_2 are regularization parameters. Using stochastic gradient descent, it is possible to update the learner feature matrix β_1 and the knowledge point feature matrix β_2 . Specifically, for each (i,j) , the update formula is as follows:

$$\begin{aligned} \beta_{1i} &\leftrightarrow \beta_{1i} + \alpha(e_{ij}\beta_{2j} - \lambda_1\beta_{1i}) \\ \beta_{2j} &\leftrightarrow \beta_{2j} + \alpha(e_{ij}\beta_{1i} - \lambda_2\beta_{2j}) \end{aligned} \quad (29)$$

where α is the learning rate and $e_{ij} = R_{ij} - \beta_1^T \beta_2$ denotes the error. Through continuous iteration, the optimal learner feature matrix β_1 and knowledge point feature matrix β_2 are updated to obtain the interaction matrix R_{ij} between learners and knowledge points, which in turn updates the implicit set of weak knowledge points Z_{ij} .

θ_k , whose gradient can be calculated by backpropagation algorithm as in Eq. (30):

$$\frac{\partial Z}{\partial \theta_k} = \frac{\partial Z}{\partial \hat{p}} \frac{\partial \hat{p}}{\partial \theta_k} \quad (30)$$

where Z is the likelihood function, $\hat{p}$ is the probability predicted by the model, and θ_k is the k th model parameter. The backpropagation algorithm can recursively compute $\frac{\partial Z}{\partial \hat{p}}$ and $\frac{\partial \hat{p}}{\partial \theta_k}$ and multiply them by the chain rule to obtain $\frac{\partial}{\partial \theta_k}$.

The model parameters can then be updated based on the gradient using gradient descent as in Eq. (31):

$$\theta_k = \theta_k - \eta \frac{\partial L}{\partial \theta_k} \quad (31)$$

where η is the learning rate and controls the size of the update step. By continuously iterating the above update process, the optimal model parameters can be obtained.

3. Experimental results of learning path recommendation

In order to verify the effectiveness of the personalized learning path recommendation model proposed in this paper, this part of the experiment will be designed in the following three aspects. First, the learners of the online learning platform are randomly divided into two groups of 60 each, one group chooses the learning path given by the personalized learning path recommendation model of this paper, called the experimental group, while the other group learns on their own in the online learning platform according to their own needs, called the non-experimental group, and statistics of the learning data of these two groups of learners are made. Compare the convergence speed of the recommendation model in this paper with that of other algorithms, and see whether the fastest

convergence speed is always maintained under different numbers of iterations. The comparison algorithms designed for the experiment are the traditional Genetic Algorithm (GA) as well as the Ant Colony Algorithm (ACO).

In order to verify the coverage of the generated knowledge point sequence, 6 learners were randomly selected and their knowledge point coverage was calculated, Table 1 is the knowledge point coverage experiment results, the coverage of the knowledge point sequence generated by the method studied in this paper can reach more than 90%, therefore, through the collaborative filtering algorithm and the correlation between the knowledge points, the knowledge point sequence that meets the learners' own needs can be generated more accurately.

Table 1. The experimental results of knowledge point coverage

Learners	A	B	C	D	E	F
Coverage	91	95	94	96	98	92

Based on the knowledge map of college English majors, this paper adopts the algorithm of this paper to generate the sequence of knowledge points. In order to verify the feasibility of the model, this paper chooses four different models and conducts comparative experiments using PRECISION and COVERAGE as indicators. These four models are Topological Ordering Algorithm (TO), Firefly Algorithm (FA), Transfer Learning Method Based on Knowledge Graph (TL-KG) and Deep Learning Method (DL). During the experiment, the knowledge point data of 200 learners are input into different models and the output results of each model are recorded. In our simulation experiment results, the data of 6 learners were randomly selected as shown in Table 2. The experimental results show that the algorithm used in this paper is superior to the other four models in terms of accuracy and coverage, and the accuracy and coverage of the knowledge point sequences generated by learner F under this paper's algorithm are 98.9% and 93.6%, which is not only the best performance among all the models, but also the best performance among all the learners.

Table 2. The comparative results of knowledge point generation algorithms

Learners	Evaluation index	This algorithm	TO	FA	TL-KG	DL
A	PRECISION/%	87.9	72	83.8	87	83.9
	COVERAGE/%	90.7	87.2	78.6	82.5	78.1
B	PRECISION/%	79.6	72	79.3	77.6	72.8
	COVERAGE/%	78.1	69.2	74.9	72	74.1
C	PRECISION/%	97.3	95.7	94.4	92.2	85.3
	COVERAGE/%	86.3	83.8	79.6	82.4	78.6
D	PRECISION/%	83.9	78.5	79.3	79.6	83.2
	COVERAGE/%	92.5	80.4	86.4	89.9	84.9
E	PRECISION/%	75.4	70.2	71.3	69.4	74.2
	COVERAGE/%	80.4	78	79.9	74.8	71.9
F	PRECISION/%	98.9	86.4	89.8	87.5	89.6
	COVERAGE/%	93.6	84.8	83	87.3	88.4

Assuming that the time can be converged to infinity to a certain extent, the algorithm in the process of finding the global optimal solution to solve the problem, not only to find the optimal solution to solve the problem, but also to make the algorithm converge to a certain stable value in solving this kind of problem, and for the algorithms with poor convergence in solving the related problems, their reference value is also relatively low. Therefore, this paper designs a large number of simulation experiments to verify the convergence of the model in this paper. For the convergence speed of the

algorithm in this paper to solve the problem of personalized learning path recommendation, this paper randomly selects the experimental data of some learners in the results of 100 simulation experiments, analyzes the iterative data of the algorithm in this paper and draws the relevant iteration curves, the experimental results of the 6 learners are shown in Figure 3-Figure 8, this paper lists the convergence of the 6 learners under the algorithm of this paper, GA and ACO, from the overall experimental results, in the algorithm convergence comparison result graph of the 6 learners, The convergence speed of this paper is significantly better than that of the traditional GA algorithm and the ACO algorithm, and in the comparison of the convergence results of the algorithm of learner E, although the convergence speed of the ACO algorithm and the GA algorithm is better than that of the proposed algorithm at first, with the increase of the number of iterations, the convergence of the proposed algorithm is significantly better than that of the ACO algorithm and the GA algorithm when the iteration is about 20 times. In summary, it can be seen that the proposed algorithm is better than the other two algorithms in solving the problem of personalized learning path recommendation.

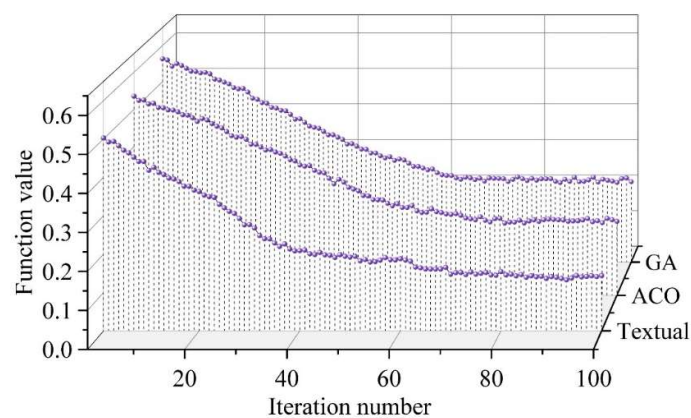


Fig. 3. The comparison of the algorithm of the learner A

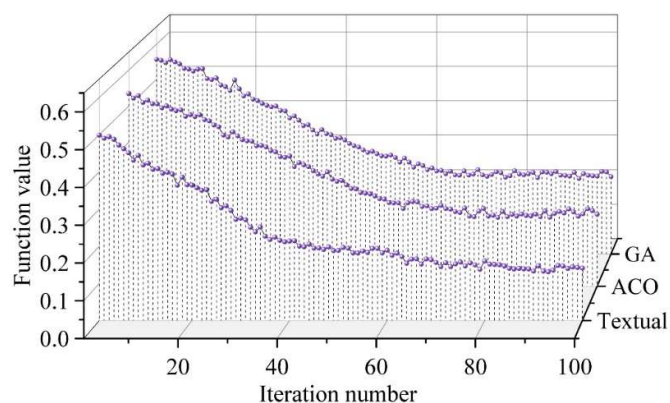


Fig. 4. The comparison of the algorithm of the learner B

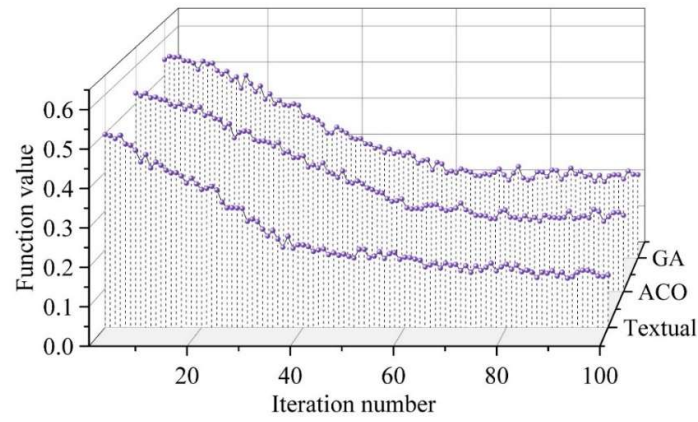


Fig. 5. The comparison of the algorithm of the learner C

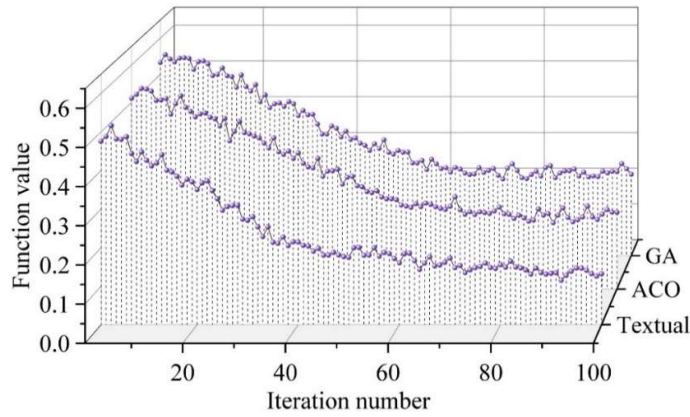


Fig. 6. The comparison of the algorithm of the learner D

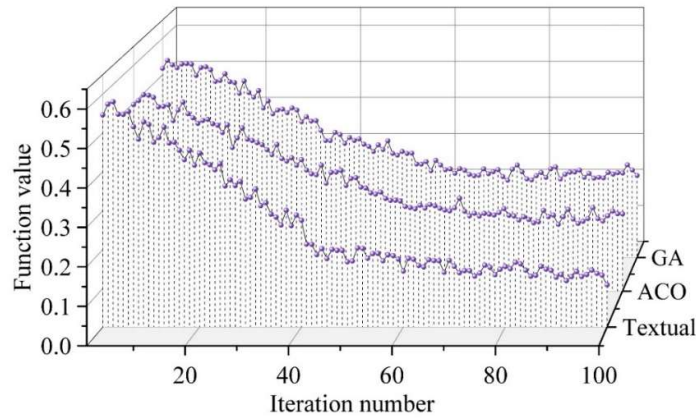


Fig. 7. The comparison of the algorithm of the learner E

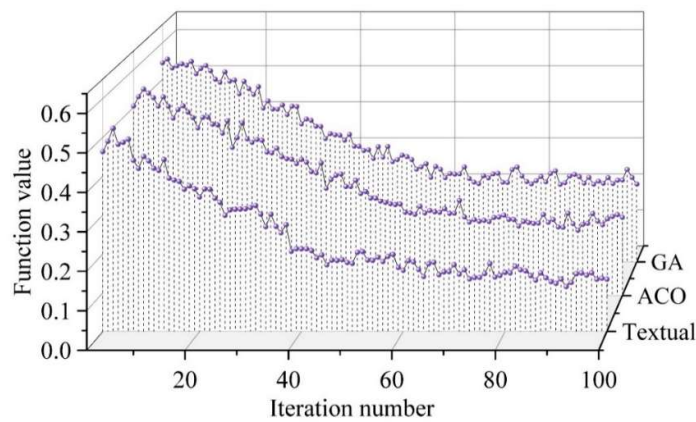


Fig. 8. The comparison of the algorithm of the learner F

Figure 9 shows the visualization results of learning path generation under different methods, Path A is a typical linear learning path of students in MOOCCubex, B is a typical circular learning path of students in MOOCCubex, and C is the learning path automatically generated by the algorithm. Path A is a typical linear learning path. The student learns all the points in sequence and does not review them during the learning process. This can easily lead to a large number of “underdeveloped”, “unlearned” and “unmastered” knowledge points in the learner's progress. As a result, this type of learning tends to be inefficient.

Path B is a typical cyclical learning path. Throughout the learning process, students will independently review their knowledge points. However, in the whole learning process, it is easy to appear that good not fully mastered knowledge points are omitted, while some mastered knowledge is repeatedly reviewed. Therefore, this type of learning is also inefficient.

In contrast to Path A and Path B, Path C is a typical learning path automatically generated for the algorithm. It accurately covers all unlearned, unmastered and insufficiently mastered knowledge points. These knowledge points are also ranked according to prerequisite relationships and difficulty levels. The learning paths generated by this dynamic learning path planning algorithm are adaptive for students. Personalized learning paths based on learning states save time and improve learning efficiency.

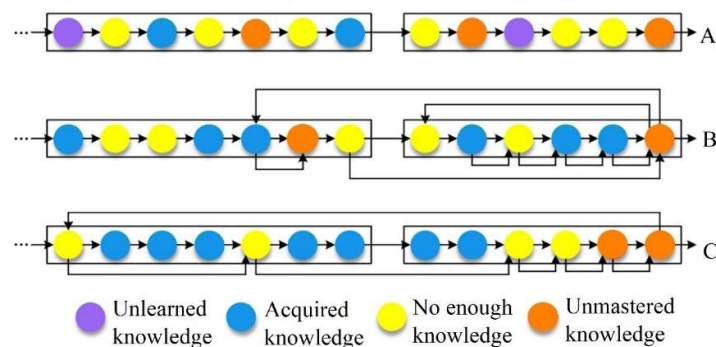


Fig. 9. Three typical learning paths

4. Analysis of the effect of model application

4.1. Statistics and analysis of average study time

At the end of the learning process, the data on students' learning time recorded in the log of the online learning platform are collected, and the average learning time of the two groups of students learning

each knowledge point can be calculated, and the statistical results are shown in Figure 10. For most of the knowledge points, the time used by the students in the experimental group is obviously shorter than that of the students in the control group, and the average learning time of the thirty learners is 82 minutes, which is much lower than that of the students in the control group. This is because when students in the experimental group study according to the recommended personalized learning path, they do not have to spend time on searching and selecting resources, so they can spend more effective time on learning, and the use of time is more efficient. However, some students are more special, this is because for some students in the experimental group, although the learning path is given, but due to their own learning ability is poor, while some students in the control group have strong learning ability, and can also find suitable learning resources. Secondly, the gap between the learning time of the students in the control group is large, while the time used by the students in the experimental group is relatively stable, this is because the students in the control group have different ability to search for resources and learning ability, and the time used to find suitable resources for themselves is different, while the experimental group is based on the resources determined by the students' learning styles and cognitive abilities, which can better meet the needs of the students, but due to the differences in the learning ability of the students, the learning time also. There is a certain gap, but compared with the control group, it is still relatively stable.

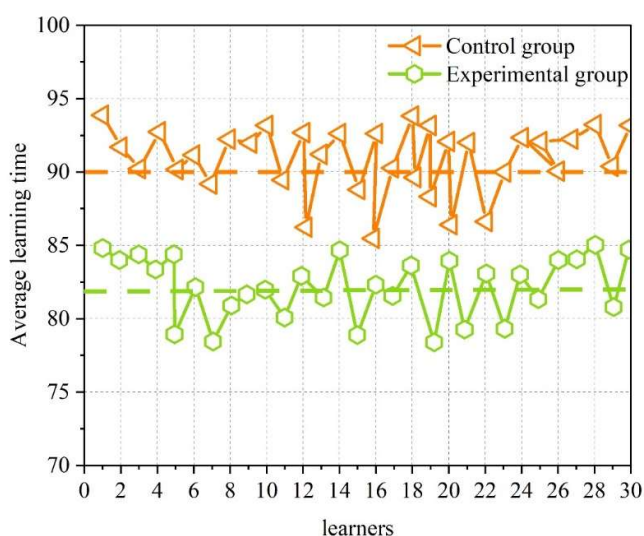


Fig. 10. Two groups of average study time comparison

In order to further verify that it is plausible that there is a significant difference between the two groups of students in the average study time, the article uses SPSS data analysis software to test the study time data, using the grouping category as the grouping variable and the study hours as the test variable, and using the independent samples t-test to compare the difference between the two groups to first determine whether the data are in line with the normal distribution, and then perform the independent samples t-test, and the results of the analysis are shown in Table 3 shown. Where the Sig value of the Leneve test of the variance equation is significantly greater than 0.10, the results of equal variances are taken, and vice versa, the results of unequal variances are taken. In terms of learning time, the Sig value of the Leneve test of the variance equation is 0.143, then the result of equal variances is taken, which means that the significance value of the difference of the t-test is 0.011 less than <0.05 , indicating that it is plausible that there is a significant difference between the students of the experimental group and those of the control group in the average learning time, which suggests that the students of the experimental group, with the support of the personalized learning paths, have alleviated the information disorientation and the cognitive load phenomenon and effectively save

learning time. Therefore, the average learning time of students learning according to personalized learning paths is significantly lower than that of traditional learning methods.

Table 3. Study time difference analysis

		Learning length	
		Let's say the variance is equal	Let's say that the variance is not equal
Levene test	F	2.219	
	Sig.	0.143	
T test	t	13.659	13.458
	df	63	60.125
	Sig.(Double side)	0.011	0.004
	Mean difference	6.768	6.526
	Standard error value	0.498	0.668
	95%CI Lower limit	5.786	5.498
	95%CI Upper limit	7.752	7.568

4.2. Analysis of Learning Achievement

After the learning activity, the pre-test data of the two groups of students were first analyzed and the results are shown in Table 4. The average score of the experimental group is slightly higher than that of the students in the control group ($72.649 > 72.546$), and the significance value of the difference (Sig. (double-test)) is 0.704, which is greater than 0.05, indicating that there is no significant difference between the pre-test data of the two groups of students, which means that the mastery of the English knowledge of the two groups of students is at the same level.

Table 4. Matching sample statistics

Group	N	Mean	Standard deviation	Standard error of mean	Sig.2
Experimental group	30	72.546	3.074	0.425	0.704
Control group	30	72.649	3.156	0.435	

The paired sample t-test was conducted on the pre-test and post-test scores of the two groups of students using SPSS, and the paired sample correlation coefficients are shown in Table 5. After learning activities in different ways, the average scores of students in the experimental group are significantly higher than the average scores of the control group, which indicates that learning according to the recommended paths can effectively improve the learning effect. Secondly, the trend of standard deviation reduction of students in the experimental group is larger, indicating that the degree of individual student discrete is smaller in the experimental group, which reflects the effectiveness of personalized learning paths to a certain extent. In addition, the significance values of students in the control group and the experimental group are both less than 0.05, indicating that the pre-test and post-test scores of students in the two groups are significantly and linearly correlated. However, the correlation coefficient of the experimental group is larger than the correlation coefficient of the control group, indicating that the pre and post-test scores of the students learning based on the recommended learning paths in this paper are more correlated, which is due to the lack of the influence of other correlation factors, and thus the correlation with the learning ability is stronger.

Table 5. Matching sample correlation coefficient

Group	N	Mean	Standard deviation	Standard error of mean	Correlation coefficient	Sig.
-------	---	------	--------------------	------------------------	-------------------------	------

Control group	Pretest	30	72.649	3.156	0.435	0.645	0.004
	Posttest	30	78.475	2.398	0.321		
Experimental group	Pretest	30	72.546	3.074	0.425	0.794	0.002
	Posttest	30	89.251	1.457	0.211		

The paired sample t-test results of the two groups of students are shown in Table 6, the difference between the mean scores of the pre-test and the post-test of the control group is -7.654, while the difference between the mean scores of the pre-test and the post-test of the experimental group is -16.487, the significance value of the differences between the two groups is less than 0.05 and the significance value of the differences of the experimental group is smaller, which indicates that the difference between the mean scores of the pre-test and the post-test of the experimental group is more significant. Therefore, it can be seen that students in the experimental group and control group have improved in terms of relevant knowledge and skills after learning through different learning paths, while students in the experimental group, with the support of the recommended personalized learning paths, reduce the time spent on searching for learning resources and devote themselves to the learning process, and secondly, according to the learning styles and knowledge structure of students, learning resources are recommended to meet the students' learning needs and knowledge abilities, students have high enthusiasm for learning, and therefore they achieve good academic results. Therefore, the learning effect of students learning according to the personalized learning path is significantly higher than that of traditional learning methods.

Table 6. Matching sample test

Group		Pair difference			t	df	Sig.2
		Mean	Standard deviation	Standard error of mean			
Control group	Pretest-Posttest	-7.654	1.225	0.195	-36.459	31	0.013
Experimental group	Pretest-Posttest	-16.487	2.897	0.354	-38.794	29	0.003

5. Conclusion

In this paper, a personalized learning path generation model based on neural network is designed to explore the dynamic planning problem of learners' learning path in the blended teaching mode of English majors.

Different personalized learning path recommendation methods are simulated and compared, and the experimental results show that this paper's algorithm outperforms the other four models in terms of accuracy and coverage, and the accuracy and coverage of the sequence of knowledge points generated by learner F under this paper's algorithm are 98.9% and 93.6%, respectively, which is the best performance. And the convergence of the collaborative filtering algorithm also achieved better results than the ACO and GA algorithms overall. The effectiveness of this paper's algorithm to solve the personalized path recommendation problem is verified.

Under the application of this dynamic learning path generation model, students can improve the efficiency of learning time utilization and find a more suitable learning path in a shorter time. The results of the paired-sample statistics of the pre-test data show that the significance value of the difference between the students in the experimental group and the control group is 0.704, which is greater than 0.05, and there is no significant difference in the learning effect. However, the paired-sample t-test results of the two groups of students show that the significance value of the difference between the learning effects of the two groups of students is less than 0.05, but the significance value of the difference in the experimental group is smaller, indicating that the difference between the mean value of the learning effects of the pre-test and post-test of the experimental group is more significant,

which shows that personalized paths to study generated according to the neural network-based dynamic learning path generation model can further promote the improvement of learning performance.

References

- [1] Zhou, Z. (2022, July). Evaluation Method of English Online and Offline Mixed Teaching Quality Based on Three-Dimensional Teaching. In *International Conference on E-Learning, E-Education, and Online Training* (pp. 549-561). Cham: Springer Nature Switzerland.
- [2] Abbott, M. L. (2019). Selecting and adapting tasks for mixed-level English as a second language classes. *TESOL Journal*, 10(1), e00386.
- [3] Asaad Hamza Sheerah, H. (2020). Using blended learning to support the teaching of English as a foreign language. *Arab World English Journal (AWEJ) Special Issue on CALL*, (6).
- [4] Triana, N. (2023). Students' Perception of Using Mixed Languages in Teaching English at SMPN 4 Buton Tengah. *Journal of Teaching of English*, 8(3), 274-283.
- [5] Sheng, L. (2020, September). Discussion on the feasible application of "Online"+"Offline" mixed learning models of English in higher vocational colleges. In *2020 International Conference on Modern Education and Information Management (ICMEIM)* (pp. 653-656). IEEE.
- [6] Kim, S. A., & Morita-Mullaney, T. (2020). When preparation matters: A mixed method study of in-service teacher preparation to serve English learners. *Mid-Western Educational Researcher*, 32(3), 231-254.
- [7] Li, G. (2021, May). A corpus-based study on the construction of online+ offline mixed mode in college English teaching. In *2021 2nd International Conference on Computers, Information Processing and Advanced Education* (pp. 1430-1433).
- [8] Yanfei, M. (2021). Online and offline mixed intelligent teaching assistant mode of english based on mobile information system. *Mobile information systems*, 2021(1), 7074629.
- [9] Ying, L. (2020). Discussion on the reform of mixed-English teaching mode of higher vocational English based on "internet+". *Journal of Contemporary Educational Research*, 4(12).
- [10] Wu, F. (2020, December). Discussion on mixed teaching mode of college English under the background of Internet. In *2020 5th International Conference on Mechanical, Control and Computer Engineering (ICMCCE)* (pp. 1659-1662). IEEE.
- [11] Chen, S., & Tian, Y. (2023). Construction of network mathematics model for college English mixed teaching. *Applied Mathematics and Nonlinear Sciences*.
- [12] Li, P., Zhang, H., & Tsai, S. B. (2021). A new online and offline blended teaching system of college English based on computer internet technology. *Mathematical Problems in Engineering*, 2021(1), 3568386.
- [13] Li, L. (2022, December). Application research of online and offline mixed teaching mode based on mobile education platform in higher vocational English teaching. In *2022 3rd International Conference on Artificial Intelligence and Education (IC-ICAIE 2022)* (pp. 13-19). Atlantis Press.
- [14] Tan, Y. (2018, December). A Study on "Online+ Offline" Mixed Teaching Mode of Public English in Higher Vocational Colleges. In *2018 8th International Conference on Management, Education and Information (MEICI 2018)* (pp. 1227-1230). Atlantis Press.
- [15] Yang, P. (2024). Study on the Influence of Mixed English Teaching on Students' Academic Performance in Higher Vocational Colleges. *Curriculum and Teaching Methodology*, 7(2), 194-199.

- [16] Halil Ibrahim Kurt & Wenxian Shen. (2024). Erratum: Finite-Time Blow-Up Prevention by Logistic Source in Parabolic-Elliptic Chemotaxis Models with Singular Sensitivity in Any Dimensional Setting. SIAM Journal on Mathematical Analysis(6),8136-8148.
- [17] Fokina D. A. & Zinchenko A. S. (2024). Selection of Collaborative Filtration Methods in Segmentation of the Customer Base. Russian Engineering Research(4),605-608.
- [18] Huang Jing & Ma Keyu. (2023). SRU-based Multi-angle Enhanced Network for Semantic Text Similarity Calculation of Big Data Language Model. International Journal of Information Technologies and Systems Approach (IJITSA)(2),1-20.