

Research on Intelligent Generation and Adaptation System of Civic Education Content Based on Machine Learning

Meihua Zhou¹, Jianliang Shen², Hua Zhang³, 

¹ Youth League Committee, Zhejiang Technical Institute of Economics, Hangzhou, Zhejiang, 310000, China

² New Product Division, Hangzhou Huaxin Mechanical and Electrical Engineering Co., Ltd., Hangzhou, Zhejiang, 310030, China

³ Youth League Committee, Zhejiang Gongshang University, Hangzhou, Zhejiang, 310000, China

ABSTRACT

This paper deeply analyzes the innovative application and intelligent upgrading steps of Artificial Intelligence Generated Content (AIGC) in Civic and Political Education. Based on metadata, we construct an automated generation model of Civics education resources, divide the meta-properties of education knowledge resources, set up a knowledge tracking model DT-BKT to obtain students' mastery of Civics knowledge, adopt personalized recommendation model to realize the high adaptability of education resources based on students' Civics learning, and combine the functions of each model to build a Civics education content intelligent generation and adaptability system. Knowledge tracking experiments show that the AUC and R2 indexes of the DT-BKT model in this paper are better than those of other comparative models, and it can better simulate the response of learners on the dataset. Facing different groups of learners is able to recommend Civics courses that meet the learners' abilities. For active learners and potential learners, the average difficulty of the recommended client layer is higher by 0.08~0.15 and 0.06~0.085 respectively, while the overall difficulty difference for inactive learners is between -0.01~0.015, and the recommended difficulty is in line with the characteristics of the learner groups.

Keywords: aigc technology, knowledge tracking, course personalized recommendation, educational resource generation, civic education

1. Introduction

Artificial Intelligence (AI), as an emerging technology, is developing rapidly and showing great potential in various fields. Among them, generative AI, as one of the important branches of AI, is not only widely used in the fields of art and literature, but also shows innovative application prospects in the field of education [1-4]. Generative AI is a technology based on machine learning, which generates

 Corresponding author.

E-mail address: 13958001401@163.com (H. Zhang).

Received 12 January 2024; Revised 05 March 2024; Accepted 25 December 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-499](https://doi.org/10.61091/jcmcc127b-499)

© 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/>).

new content, such as images, music, text, etc., by learning a large amount of data [5-7]. Generative AI can be categorized into rule-based approaches and neural network-based approaches. Rule-based methods rely on manually formulated rules to generate content, while neural network-based methods generate content by training neural network models, the most common method of which is the use of generative adversarial networks [8-11]. Generative artificial intelligence is of great significance in the generation of content for civic and political education in colleges and universities.

Civic and political education refers to cultivating students' correct values, social outlook, political awareness, moral character, and, sense of social responsibility through education and cultivation, as well as improving their critical thinking and problem-solving abilities [12-15], shaping their personality qualities and moral literacy, so as to make them become socialist builders and successors in the new era who have ideals, responsibilities and sense of responsibility [16-18]. Civic education is an important part of modern education and one of the most important ways to cultivate socialist builders and successors with all-round development of morality, intelligence, physicality, aesthetics and labor. In contemporary society, the importance of Civic and Political Education is emphasized [19-22].

In this paper, the intelligent upgrading steps of ideological and political education based on AIGC technology are taken as the framework of the whole text, which is based on the intelligent upgrading steps of "automatic generation of teaching resources", "student learning situation analysis and teaching evaluation", "personalized course recommendation" and "intelligent learning platform", and finally builds a model and finally builds an intelligent generation and adaptation system for ideological and political education content. The M-S-S educational resources aggregation model based on metadata, social annotation and social network analysis is proposed. Taking advantage of the relational hierarchical and framework normative features of metadata to enhance the precision and comprehensiveness of learning resources characterization, the meta-attributes of Civic Resources and Knowledge are divided into three categories, namely, descriptive information, attributional information, and subsidiary information. Based on the standard Bayesian knowledge tracking model, the difficulty coefficients of exercises are introduced and incorporated into the performance states of the BKT model to improve the performance of model prediction. Combined with the automated generation of ideological education resources, knowledge tracking, and course personalized recommendation model for functional development, the front-end and back-end separation architecture, MVC layered mode to build the ideological education content intelligent generation and adaptation system. In order to test the knowledge tracking performance of the model and system in this paper, knowledge tracking experiments are carried out, and at the same time, a large number of Civics course resources of an online education platform constitute the experimental dataset, so as to carry out intelligent education content generation and adaptation experiments, and to validate the reasonableness of the personalized recommendation function of Civics courses.

2. Innovative Application of AIGC Technology in Civic and Political Education

Digital transformation based on machine learning is not only an important task of the times for higher vocational colleges and universities, but also a necessary way to modernize the Civic and Political Education. In this context, how to promote the digital transformation of Civic and Political education by using Artificial Intelligence Generated Content technology (AIGC technology) has become an important issue for higher education institutions to realize high-quality development [23]. Through AIGC technology, the steps of intelligent upgrading of Civic and Political Education in higher education institutions are sequentially as follows.

2.1. Automated generation of teaching resources

AIGC technology is used as the basis for building the model of Civic and Political Education Resources. Using AIGC technology, teaching courseware, practice questions and vocational skills training tutorials can be automatically generated. AIGC tools can understand and imitate the way of human creation of content, and automatically generate corresponding teaching materials according to the syllabus and knowledge points provided by the teachers, which can effectively reduce the workload of the teachers and improve the efficiency of the generation of teaching resources.

2.2. Learning situation analysis and teaching assessment

Utilizing AIGC technology and using knowledge tracking as a means to analyze the learning situation and teaching assessment of students' Civic Education. Automatically collect and analyze students' activity data on the online learning platform, such as the time spent watching courses, the speed of completing assignments, the frequency of participating in discussions, etc., so as to assess students' learning motivation and concentration. Accurately assess students' mastery of each knowledge point, identify learning difficulties and blind spots, and provide teachers with targeted teaching suggestions.

2.3. Personalized course recommendation

Relying on the knowledge tracking model for real-time tracking of students' Civic and Political knowledge learning status, the personalized course recommendation model is constructed by combining with AIGC technology to provide students with personalized learning paths and recommended course resources, and generate personalized learning content.

2.4. Intelligent Learning Platform

Integrate the above-mentioned implementation methods of "automatic generation of teaching resources", "student learning situation analysis and teaching evaluation" and "personalized course recommendation", set up an intelligent generation and adaptation system for ideological and political education content, build a data collection and processing module, collect and clean students' learning behavior, grades, interests and preferences, generate personalized learning content, and recommend learning resources for related courses according to learning progress and interest preferences.

In the following research, the intelligent upgrading steps of ideological and political education will be taken as the framework of the whole text, and the intelligent generation and adaptation of ideological and political education content will be analyzed in depth.

3. Automated Generation Model of Civic and Political Education Resources

The first and foremost foundation of the intelligent generation and adaptation of Civic and Political Education content is to realize the automated generation of Civic and Political Education resources and to provide rich and massive Civic and Political Education content.

3.1. Modeling

This study proposes an M-S-S educational resource aggregation model based on metadata, social annotation and social network analysis. The model consists of three layers, the bottom layer is the metadata-based educational resource organization framework, which overall limits the scope of annotation on resources; the middle layer is the social annotation layer, where users construct users' understanding and feedback on resources by tagging them. The top layer is the social network analysis layer, which is mainly used to analyze the social annotation results of the middle layer in terms of neutrality, sexuality, correlation, and aggregation, etc. The M-S-S educational resource aggregation model combines the resource characterization method of metadata, gives play to the relational

hierarchical and framework normative features of metadata, and embeds the social annotation function into the meta-numerical framework sub-structures to characterize and describe the educational resources, so as to improve the accuracy, accuracy, and accuracy of learning resource characterization and to improve the quality of learning resources. The accuracy and comprehensiveness of learning resources characterization. The multi-dimensional description of social annotation increases the flexibility of categorization, facilitates resource investigation and sharing, and provides technical support for active service of information resources. Based on the rich and diverse labeling results, the clustering analysis of social network analysis can be used to associate synonym labels and related labels to form a label cloud, thus converging the loose labeling structure and broad labeling range into a more cohesive structure.

3.2. Organization of educational resources

In this study, educational resource representations are considered to be divided into 3 main categories, descriptive information, attributional information, and collateral information. The first category is mainly divided into two categories, educational and technical descriptions. Educational describes the content of educational resources, such as number, name, content introduction, applicable objects, etc. Technical describes the technical indexes and technical characteristics of educational resources, including media format, media form, file size, clarity, playback time, etc. The second category mainly reflects the information of the subject, applicable objects, editors, etc. The second category mainly reflects information such as the subject to which it belongs, the applicable object, and the editor. The third category mainly reflects the creation time, update time and other information. The metadata attributes of educational resources are shown in Table 1. For the “keyword” attribute, social labeling can be used to allow learners to provide feedback and evaluation of the resource after it has been used. The attributes “type of interaction” and “degree of interaction” can also be socially labeled by the learner to provide descriptions and feedback. In the attribute of “suggested learning time”, it can be analyzed statistically based on the actual learning time of the learner in the learning process, and combined with the results of social labeling, and then weighted with the original value of “suggested learning time” to obtain the average. Learners can also describe and label the “difficulty” attribute. The use of social tags to represent educational resources improves understanding and collaboration among users, and allows resources to be labeled into categories that are most acceptable to the user community.

Table 1. Metadata attribute of education resources

Categories		Attribute
Description information	Educational attribute description	Numbering
		Name
		Content profile
		Keywords
		Resource type
		Interaction type
		Degree of interaction
		Cognitive type
		Recommended duration
		Difficulty
		Media form
		Media format
		File size
		Articulation
Attribution information	Attribute element	Playback length
		Subject of the naked course
		Applicable object
Attached information	Accessory element	Author
		Creation time
		Update time
		Version number

4. Knowledge Tracking Model for Civic Education

4.1. Fundamentals of the BKT model

To better explain the learner's learning process, the BKT model establishes four parameters for each knowledge point, $P(L_0)$, $P(T)$, $P(G)$ number $P(S)$.

1) Define the probability of answering question t correctly and incorrectly

First, based on the probability distribution table, the probability that learner u correctly answers question t involving knowledge point k is defined as the sum of the probability that the learner does not make a mistake if he/she has mastered knowledge point k and the probability that he/she guesses correctly if he/she has not mastered knowledge point 4 as shown in Equation (1):

$$P(O_t = 1)_u^k = P(L_t)_u^k \times (1 - P(S)^k) + (1 - P(L_t)_u^k) \times P(G)^k \quad (1)$$

where k denotes the knowledge point, $P(L_t)_u^k$ denotes the probability that learner u grasps knowledge point k when answering question t , $P(S)^k$ denotes the probability that the learner misses on knowledge point k , and $P(G)^k$ denotes the probability that the learner guesses on knowledge point k :

$$P(O_t = 0)_u^k = P(L_t)_u^k \times P(S)^k + (1 - P(L_t)_u^k) \times (1 - P(G)^k) \quad (2)$$

2) Calculate the probability of mastering knowledge point k under different answering scenarios respectively

According to Bayes' theorem, it can be deduced that the probability of learner u mastering knowledge point k in the case of answering question t correctly is [24]:

$$P(L_t | O_t = 1)_u^k = \frac{P(L_t)_u^k \times (1 - P(S)^k)}{P(O_t = 1)_u^k} \quad (3)$$

Similarly, the probability that learner u grasps knowledge point k in the case of incorrectly answering question t is:

$$P(L_t | O_t = 0)_u^k = \frac{P(L_t)_u^k \times P(S)^k}{P(O_t = 0)_u^k} \quad (4)$$

3) Calculate the a priori probability of learning and the current probability of knowledge acquisition

The a priori probability that a learner has mastered a knowledge point before answering question $t+1$ is the sum of the probability of mastering a knowledge point when answering question t correctly and the probability of mastering a knowledge point when answering question t incorrectly, as shown in equation (5).

$$P(L_t)_u^k = P(L_t | O_t = 1)_u^k + P(L_t | O_t = 0)_u^k \quad (5)$$

It is then deduced that the probability of the learner mastering the knowledge point when answering question $t+1$ is the probability of knowledge mastery when answering question t versus the probability of learning without mastering the knowledge point, as shown in equation (6).

$$P(L_{t+1})_u^k = P(L_t)_u^k + (1 - P(L_t)_u^k) \times P(T)^k \quad (6)$$

4) Predicting the probability of correctly answering the next question

The probability of answering Question $t+1$ correctly is the sum of the probability that the learner did not miss if he/she has mastered the point and the probability that he/she guessed correctly if he/she has not mastered the point, as shown in Equation (7).

$$P(O_{t+1} = 1)_u^k = P(L_{t+1})_u^k \times (1 - P(S)^k) + (1 - P(L_{t+1})_u^k) \times P(G)^k \quad (7)$$

4.2. Presentation of the DT-BKT model

In this section, we consider starting from the level of exercises, obtaining the difficulty of exercises based on learners' previous answers, incorporating the difficulty of exercises into the performance state of the BKT model, constructing the Bayesian Knowledge Tracking based on Exercise Difficulty (D-BKT) model, and exploring the effect of the difficulty of exercises on the prediction of knowledge tracking [25].

4.2.1. Method of calibrating the difficulty of exercises. For the knowledge tracking model that introduces the difficulty of exercises, determining the difficulty coefficient of the exercises has a great impact on the accuracy of the prediction results, which directly affects the accuracy of the assessment of learners' learning effects. At present, the two commonly used methods for difficulty assessment of exercises are pre-difficulty calibration and post-difficulty calibration.

1) Pre-difficulty calibration. Pre-difficulty calibration refers to estimating the difficulty of an exercise based on the exercise itself, before the test or during the preparation of the test questions, which mainly includes difficulty calibration based on experience, difficulty calibration based on cognitive theory and difficulty calibration based on mathematical models.

2) Post-hoc difficulty calibration. Post-facto difficulty calibration refers to the use of statistical methods to estimate the difficulty of the exercises based on the students' performance after the test, which is the simplest and easiest method at present.

(1) Calculation using pass rate. For 0-1 scoring exercises, the ratio of the number of correctly answered items to the total number of participants in the quiz is generally used as an indicator, i.e., the passing rate of the exercises is used as an indicator of difficulty, and the mathematical expression is as follows:

$$P = \frac{R}{N} \quad (8)$$

where, P represents the difficulty of the exercise, R represents the number of people who answered the exercise correctly, and N represents the total number of people who took the quiz.

(2) Calculation using average score. For multilevel scoring exercises, the difficulty is defined as the ratio of the average score of all students participating in the quiz to the full score of the exercise, with the following mathematical expression:

$$P = \frac{\bar{X}}{W} \quad (9)$$

where P denotes the difficulty of the exercise, $\bar{X}$ denotes the average score for the exercise, and W denotes the full mark value for the exercise.

First, the number of students who incorrectly answered Exercise i was counted:

$$Incorrect_i = \sum_{j=1}^S count(o_{ji} = 0) \quad (10)$$

where i denotes the exercise serial number, j denotes the student serial number, S denotes the total number of students, o_{ij} denotes the behavior of students j in answering on exercise i , and $o_{ij} = 0$ denotes answering errors.

Next, the total number of students who made response behaviors (both correct responses and incorrect responses) on Exercise i is counted:

$$Total_i = \sum_{j=1}^S count(o_{ji}) \quad (11)$$

Finally, the ratio of the number of people who incorrectly answered the exercises to the total number of people who took the quiz, D_i , was calculated and used as the difficulty factor with the following formula:

$$D_i = \frac{Incorrect_i}{Total_i} \quad (12)$$

Although the public dataset is more convenient to obtain, the data are mixed, and it is not possible to ensure that the number of each exercise is up to standard, and it is possible that the amount of data in some exercises is too small, which in turn leads to the calculated difficulty coefficient of 0 or 1, which does not reflect the real difficulty of the exercises. Considering that this is a dichotomous problem, this study uses the Laplace correction method for correction, as shown in Equation (13):

$$D_i = \frac{Incorrect_i + 1}{Total_i + 2} \quad (13)$$

4.2.2. Model structure and probability formulas. The D-BKT model adds a new exercise node D to the standard BKT model, which is used to introduce the difficulty coefficient of the exercise into the knowledge tracking model. At this time, the parameter update formula corresponding to the D-BKT model is:

$$P(G)_t^k = \frac{P(G)^k}{1 + D_t} \quad (14)$$

$$P(S)_t^k = P(S)^k \times (1 + D_t) \quad (15)$$

For Knowledge Point k , the conditional probability formula corresponding to the D-BKT model is:

1) When the answer is shown to be correct, two cases are considered: the first one, the student has mastered the knowledge point and does not make mistakes; the second one, the student has not mastered the knowledge point but is able to answer the exercise correctly by guessing. Similarly, when the answer is shown to be incorrect, there are two cases: the first, the student has mastered the point but made a mistake in answering the question; the second, the student has not mastered the point and did not guess the right answer. The specific formula is as follows:

$$P(O_t = 1)_u^k = P(L_t)_u^k \times (1 - P(S)_t^k) + (1 - P(L_t)_u^k) \times P(G)_t^k \quad (16)$$

$$P(O_t = 0)_u^k = P(L_t)_u^k \times P(S)_t^k + (1 - P(L_t)_u^k) \times (1 - P(G)_t^k) \quad (17)$$

where $P(O_t = 1)_u^k$ denotes the probability that for question t under knowledge point k , student u answers correctly; $P(O_t = 0)_u^k$ denotes the probability that for question t under knowledge point k , student u answers incorrectly; $P(S)_t^k$ denotes the probability that a student's mistake on question t on knowledge point k leads to a wrong answer; $P(G)_t^k$ denotes the probability that a student's guess on question t on knowledge point k is right and leads to a correct answer

2) Update the learning prior probability based on the sequence of students' answers, which can be expressed by the following formula:

$$P(L_t | O_t = 1)_u^k = \frac{P(L_t)_u^k \times (1 - P(S)_t^k)}{P(L_t)_u^k \times (1 - P(S)_t^k) + (1 - P(L_t)_u^k) \times P(G)_t^k} \quad (18)$$

$$P(L_t | O_t = 0)_u^k = \frac{P(L_t)_u^k \times P(S)_t^k}{P(L_t)_u^k \times P(S)_t^k + (1 - P(L_t)_u^k) \times (1 - P(G)_t^k)} \quad (19)$$

$$P(L_t)_u^k = P(L_t | O_t = 1)_u^k + P(L_t | O_t = 0)_u^k \quad (20)$$

where $P(L_t | O_t = 1)_u^k$ denotes the probability that student u has mastered knowledge point k for question t under knowledge point k on the basis of correct answers; $P(L_t | O_t = 0)_u^k$ denotes the probability that student u has mastered knowledge point k for question t under knowledge point k on the basis of incorrect answers.

3) The formula for updating the student's knowledge state for the next question is as follows:

$$P(L_{t+1})_u^k = P(L_t)_u^k + (1 - P(L_t)_u^k) \times P(T)^k \quad (21)$$

where $P(T)^k$ denotes the probability that a student changes from a state of no mastery to a state of mastery for knowledge point k , through learning.

4) The formula for the predicted probability that a student will answer the next question correctly is as follows:

$$P(O_{t+1})_u^k = P(L_{t+1})_u^k \times (1 - P(S)_{t+1}^k) + (1 - P(L_{t+1})_u^k) \times P(G)_{t+1}^k \quad (22)$$

In the T-BKT model, students' learning ability can be measured by the difference between the correct rate and the error rate of knowledge points. In this study, the Laplace correction method was used to correct the correct and error rates, and considering that this is a binary classification problem, the corrected formula is shown below:

$$R_k = \frac{Count_{o=1}^k + 1}{Count_{o=0}^k + Count_{o=1}^k + 2} \quad (23)$$

$$W_k = \frac{Count_{o=0}^k + 1}{Count_{o=0}^k + Count_{o=1}^k + 2} \quad (24)$$

$$T^k = \begin{cases} R_k - W_k, R_k \geq W_k \\ 0.01, R_k < W_k \end{cases} \quad (25)$$

After obtaining the learner's learning ability about the knowledge point, it is initialized by assigning it to parameter $P(T)$ with the following formula:

$$P(T) = (T^1, T^2, \dots, T^k) \quad (26)$$

5. Personalized Recommendation Model for Civic and Political Education Courses

5.1. SAE-CRD Course Recommendation Model

The SAE-CRD course recommendation model proposed in this subsection is constructed based on the autoencoder neural network structure. The autoencoder neural network structure mainly consists of encoder and decoder. The encoder maps the input data to a new vector space through the neuron's activation function thus obtaining a new feature representation vector, while the decoder computes the feature representation vector through the neuron's activation function thus obtaining the output. The SAE-CRD course recommendation model improves the encoder and decoder respectively, i.e., Self-Attention Mechanism Encoder (SAE), Course Relevance Decoder (CRD).

1) Self-attention mechanism encoder. In a self-attentive mechanism encoder, its input is the learner's initial rating matrix initialized from the learner's course learning record. Suppose the initial course rating matrix $X_s \in \mathbb{R}^N$ of learner s is obtained and the course learning record $K_s = \{k_1, \dots, k_n\}$ of the learner is given, where k_i denotes the course index.

The embedding matrix weights W_1 of the course are first obtained using an autoencoder neural network, and the acquisition of this embedding matrix is similar to the principle of acquiring the word embedding matrix of word2vec model, as shown in Eqs. (27), (28), and (29).

$$y_i = f_1(W_1 x_i + b_1) \quad (27)$$

where f_1 , W_1 and b_1 are the activation function, weight matrix and offset vector of the encoder, respectively, and y_i is the representation vector of the input data mapped to the hidden space.

$$\hat{x}_i = f_2(W_2 y_i + b_2) \quad (28)$$

where f_2 , W_2 and b_2 are the activation function, weight matrix and offset vector of the decoder, respectively, and $\hat{x}_i$ is the representation vector on the implicit space.

$$Loss = \sum_{i=1}^M \|x_i - \hat{x}_i\|_2^2 \quad (29)$$

where $Loss$ is the loss value of the scoring matrix reconstruction.

Then the embedding vectors of learner s 's historical study courses are selected from the course embedding matrix W_1 , thus composing the learner's implicit vector representation, which is computed as shown in Eq. (30):

$$W_1[K_s] = (W_1^{k_1}, \dots, W_1^{k_n}) \quad (30)$$

where $W_1 \in R^{H \times N}$, $[\]$ denote the sub-matrices of the weight matrix W_1 , which is a representation of the slices of the course embedding matrix over the corresponding history courses of the learner, and $W_1^{k_i}$ denotes the embedding vector of the course with index k_i .

Assuming that the dimension value of the attention obtained from the embedding layer is denoted by H_a , the self-attention mechanism is computed as shown in Eqs. (31), (32):

$$A_s = \text{soft max}(\tanh(W_a W_1[K_s]) + b_1) \quad (31)$$

where $A_s \in R^{H_a \times n}$ is the importance scoring matrix, where each row in this scoring matrix indicates the importance of a particular aspect of the learner's preference on different courses, and each column indicates the importance of the learner's preference on each course:

$$Y_1^s = A_s \cdot [K_s]^T \quad (32)$$

where, $Y_1^s \in R^{H_a \times H}$ denotes the learner's preference matrix also known as the learner's representation matrix.

2) Course relevance decoder. In the course relevance decoder, the input is the learner's representation matrix obtained from the encoder Y_1^s . In order to facilitate the subsequent decoding operation, a merging layer is used to merge Y_1^s from a multidimensional vector matrix to a one-dimensional vector representation, which is computed as shown in Equation (33):

$$y_1^s = f_1(Y_1^{sT} w_t + b_t) \quad (33)$$

where f_1 is the activation function of the merged layer, w_t and b_t are the weights and bias of the merged layer, respectively:

$$\begin{cases} y_2^s = f_2(W_2 y_1^s + b_2) \\ y_3^s = f_3(W_3 y_2^s + b_3) \end{cases} \quad (34)$$

where, W_2 and W_3 are the weights of the bottleneck layer, and b_2 and b_3 are the biases of the bottleneck layer, respectively.

Since the learner will be influenced by the learned courses in the course of the course, not only the learner's learning interest but also the correlation between the unlearned courses and the learned courses should be taken into account when recommending courses to the learner. In this decoder, the correlation between different courses will be obtained by using the dot product of the embedding

vectors of the learner's studied courses and the embedding vectors of the unlearned courses, which is computed as shown in Equation (35):

$$L_s = W_4 \cdot W_1 [K_s] \quad (35)$$

where $L_s \in R^{N \times n}$ denotes the influence between different courses for learner s .

When using AMSLSTM to extract the correlation between courses, it is necessary to convert the description text of the courses into word vectors first, and then use LSTM for text training. That is, a sequence of word vectors of course description statements $x_1, \dots, x_t$ is passed into the LSTM neural network, and the representation vectors of course descriptions are obtained by summation through the implicit states h_t of the LSTM. Specifically, when training with LSTM, each word vector is used as a time-step input, and the extracted hidden vectors h_t are computed as shown in Eq. (36):

$$h_t = \tanh(W \cdot x_t + U \cdot h_{t-1}) \quad (36)$$

where x_t is the input of time step t , W and U denote the weighting matrix of the input and the implied state vector of time step $t-1$ to the implied state vector of that time step, respectively.

The semantic information in the course description text is better extracted using the attention mechanism, which is calculated as shown in Eqs. (37), (38):

$$a_c = \text{soft max}(\tanh(W_l \cdot h_t + b_l)) \quad (37)$$

where W_l and b_l are the weights and biases of the attention mechanism, respectively, and a_c is the attention weight of the hidden state h_t :

$$h_i = \sum_t a_c \cdot h_t \quad (38)$$

where h_t is the implied state at time step t and h_i is the representation vector of the i th course.

The representation vectors of different courses are computed using the Manhattan distance to measure the relevance of their courses, and the relevance measure of courses numbered 1 and 2, for example, is shown in Equation (39):

$$D[K_s] = \exp(-\|h_1 - h_2\|) \quad (39)$$

where, h_1 and h_2 denote the course numbered 1 and 2 respectively i.e. the output of the upper and lower network structures in the graph. The relevance is mapped between 0 and 1 by index calculation.

The extracted relevance metric $D[K_s]$ is incorporated into the previous course relevance matrix L_s , which is calculated as shown in (40):

$$L_s = (W_4 \cdot W_1 [K_s]) \otimes D[K_s] \quad (40)$$

where $\otimes$ is the product of the elements.

Summing each row of L_s yields the influence of other courses on the course $l_s \in R^N$, which is calculated as shown in (41):

$$l_s = \sum_{j=1}^n L_s^{(j)} \quad (41)$$

Finally, the reconstruction of the course rating matrix is obtained based on the learner's preference, and the relevance of the courses, where the reconstruction of learner s 's rating for each course is calculated as shown in (42):

$$\hat{x}_s = f_4(W_4 \cdot y_3^s + l_s + b_4) \quad (42)$$

where $W_4 \cdot y_3^s$ is the preference profile representation of learner s , l_s is the course relevance for learner s , and b_4 is the bias.

5.2. Model Training

In the process of model training, we take the courses that students have enrolled in as positive examples, and the unlearned courses as negative examples, and use the weighting method to more clearly distinguish between the unlearned courses and the unlearned courses of the learners, in which a uniform weight is set for the negative examples, and for the positive examples, the more frequently the students learn, the larger the corresponding weights are. In this way, the confidence matrix of learner preference C is obtained, and its calculation process is shown in Equation (43):

$$C_{s,c} = \begin{cases} 1 + \alpha \log \left(1 + \frac{f_{s,c}}{\varepsilon} \right) \\ 1 \end{cases} \quad (43)$$

where α and ε denote the weight parameters of the confidence.

When the confidence matrix is embedded, the objective function of the model is shown in Equation (44):

$$Loss = \sum_{s=1}^M \sum_{c=1}^N \left\| C \otimes (X_{s,c} - \hat{X}_{s,c}) \right\|_2^2 \quad (44)$$

In order to be able to train the model better, a regularization term is added to the objective function. The objective function after adding the regularization term is shown in equation (45):

$$L = Loss + \lambda \left(\|W_i\|_2^2 + \|W_a\|_2^2 + \|W_t\|_2^2 \right) \quad (45)$$

where, λ is the regularization parameter, W_i is the weight matrix, W_a is the learning parameter in the attention layer, and W_t is the learning parameter in the merging layer. The addition of the regularization term not only makes the model have a better generalization effect, but also can effectively mitigate the noisy data.

6. Intelligent Generation and Adaptation System for Civic and Political Education Content

This chapter will integrate the automated generation model of civic education resources, the civic education knowledge tracking model and the personalized recommendation model of civic education courses proposed above, reasonably plan the system functional modules, realize the intelligent generation of civic education contents and resources, reasonably carry out the personalized recommendation of civic education courses according to the students' civic education knowledge learning, and complete the direct adaptation of civic education contents and resources to the students.

6.1. Overall system architecture

The system design adopts mainstream B/S architecture, the back-end development mainly uses SpringBoot framework, the front-end page mainly uses Vue.js framework and ElementUI. Data persistence using Mybatis, Redis and other technology components, Mybatis as a widely used open source ORM framework, through object-oriented syntax to achieve convenient and fast data persistence and other operations. Redis, as a commonly used in-memory database, has faster read and write speeds compared to traditional disk-stored databases, and can often be used to improve

access speeds, reduce database load, and improve application performance. The system design uses RestfulApi to standardize the interaction between the front and back ends.

The system adopts front-end and back-end separation architecture, MVC layered model, mainly including the underlying operating environment, data layer, business layer, and view layer. The specific architecture is shown in Figure 1.

The details of each layer from bottom to top are described below.

- 1) Running environment. The basic platform is the environment on which the system runs normally, this system uses Linux development environment, the system uses Java language and SpringBoot to develop and build.
- 2) Database layer. Database layer is mainly for the system application data and user data persistence operations, the system mainly uses Mysql as the data storage database, Redis do system cache database.
- 3) Data Layer. This layer is mainly for the system to provide business data read and write operations, the system mainly uses Mybatis to realize the system business data persistence and CRUD operations. As well as the use of Redis to reduce the system business database load.
- 4) Application Layer. The business layer mainly realizes the business functions of the system, including the student side and the teacher side. The system functions are realized through modular design, and decoupling design is carried out between different modules to reduce the dependency between modules and improve independence.
- 5) Support Layer. Through HTTP request, it realizes communication between the front and back ends, interface call, cross-domain access and call of business functions.
- 6) View Layer. This layer is the display page of the system, the user through the page for the operation of the system to use, mainly using Vue.js framework and ElementUI development and design.

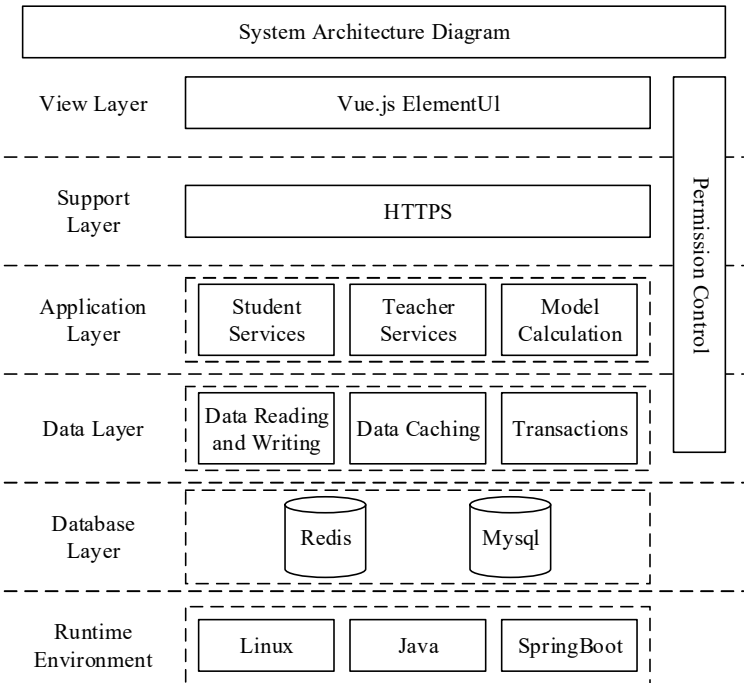


Fig. 1. System architecture

6.2. Functional Module Design

This system mainly includes two modules: the teacher's side and the student's side. The student's side mainly includes functions such as registration and login, course learning, test answering and corresponding recommended exercises according to the student's answers. On the teacher's side, it mainly includes the functions of student management, knowledge point management, test question management and so on.

7. Civic Education Knowledge Tracking Experiment

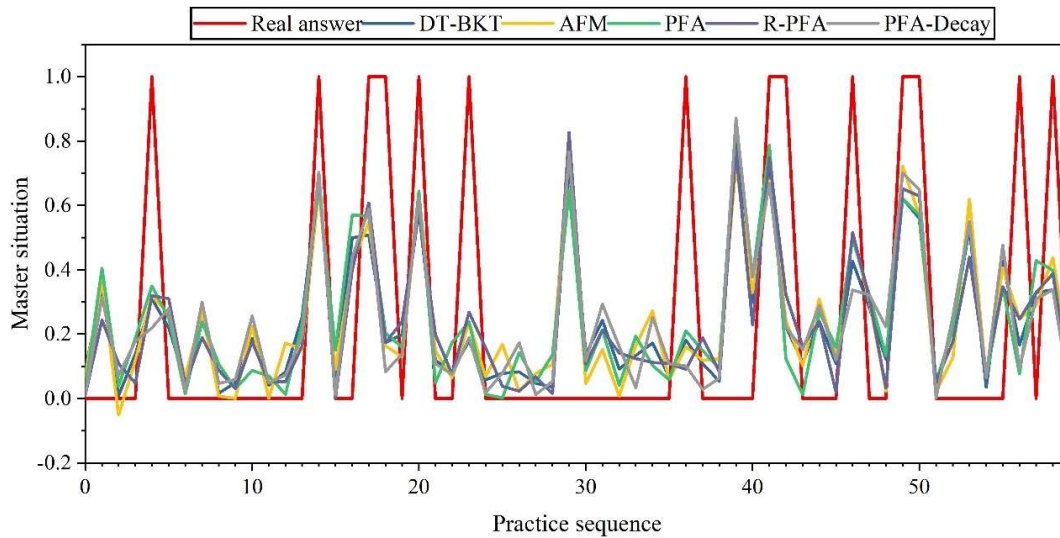
The content of this chapter is to conduct knowledge tracking experiments on DT-BKT, the knowledge tracking model of Civic Education constructed in this paper. This time, three datasets, Fraction, Tone and Cloze, are selected for knowledge tracking modeling, and AFM, standard BKT model, D-BKT model, T-BKT model, and DT-BKT model of this paper are established, and the modeling results are specifically shown in Table 2. It can be seen that the AUC and R2 metrics of the DT-BKT model proposed in this paper are better than those of the AFM model, the standard BKT model, and the extended D-BKT model and the T-BKT model of the BKT model on the public dataset. This result indicates that the DT-BKT model performs better than existing models in predicting students' future performance in civic education.

Table 2. Modeling results

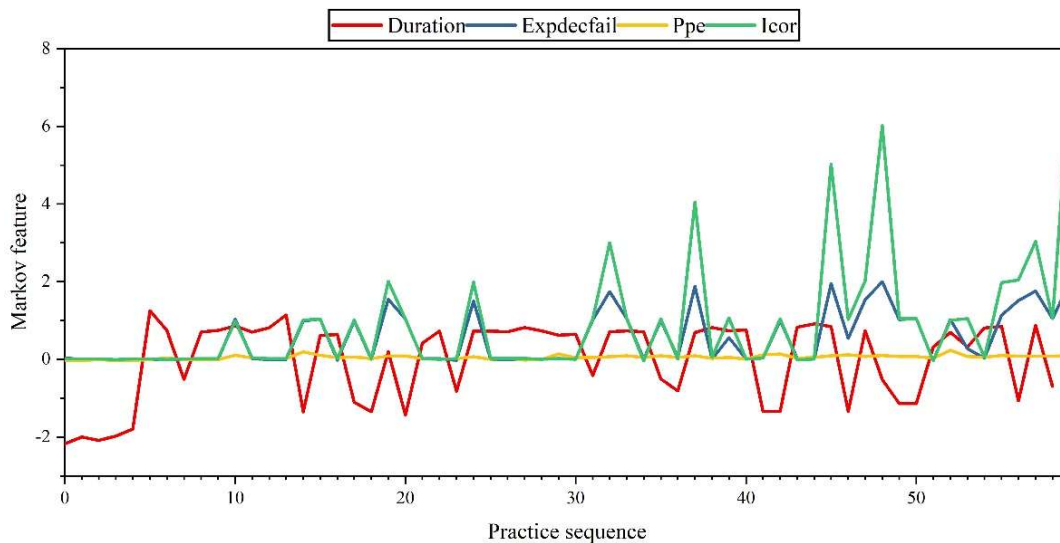
Data set	Index	AFM	BKT	D-BKT	T-BKT	DT-BKT
Fraction	Precision	0.847	0.916	0.863	0.873	0.861
	Recall	0.972	0.977	0.97	0.963	0.973
	F1	0.914	0.913	0.907	0.923	0.903
	AUC	0.758	0.778	0.789	0.792	0.833
	R ²	0.159	0.187	0.211	0.219	0.256
Tone	Precision	0.841	0.835	0.844	0.832	0.843
	Recall	0.983	0.978	0.96	0.956	0.956
	F1	0.906	0.891	0.901	0.896	0.909
	AUC	0.804	0.809	0.816	0.826	0.827
	R ²	0.212	0.219	0.223	0.212	0.227
Cloze	Precision	0.772	0.772	0.78	0.789	0.802
	Recall	0.774	0.762	0.757	0.767	0.767
	F1	0.77	0.763	0.766	0.785	0.786
	AUC	0.853	0.861	0.847	0.85	0.855
	R ²	0.291	0.308	0.316	0.328	0.344

To further explain the changes in learners' abilities during the Civics learning process, learners' correct rates on each knowledge point can be visualized. In this subsection, the answer sequence of a student in the Cloze dataset is randomly selected to visualize the changes in the Markov Blanket feature of the learner, the changes in the real answer response and the changes in the correct rate predicted by the model. The change process of learner's correct rate and the change process of Markov blanket features are shown in Figure 2. Figures (a) and (b) correspond to the change of learners' correct rate and the change of Markov blanket features in turn. The horizontal coordinate of the former is the learner's practice sequence, and the vertical coordinate is the learner's real response and the correct rate predicted by the model, respectively, while the horizontal coordinate of the latter is the learner's practice sequence, and the vertical coordinate is the Markov Blanket feature values, respectively. It can be seen that the variations of Markov blanket features on the Cloze dataset are able to model the learners' answer responses better. The learners' correctness rate increases with the

increase of the number of practice sessions. When the learners first enter the system, the answer time is shorter and the correct rate is lower, which may be due to the learners' chaotic behavior or unfamiliarity with the knowledge points at the beginning, but with the increase of the number of practice sessions, the learners gradually adapt to the system, and the answer time increases and the correct rate gradually improves.



(a) Changes in learners' accuracy



(b) Changes in the characteristics of the markov carpet

Fig. 2. Learners' accuracy and markov carpet

8. Experiments in Intelligent Educational Content Generation and Adaptation

In this chapter, we will carry out an intelligent educational content generation and adaptation experiment to evaluate and analyze the performance of the personalized recommendation model of the Civics education course and the intelligent Civics education content generation and adaptation system in this paper. The dataset for this experiment comes from an online education platform that can provide a large number of Civics course resources. There are 93 educational courses in the dataset, and each course corresponds to an accompanying quiz to understand the learners' understanding of the course content.

Learners are categorized by identity labels based on two indicators: learners' ability and the number of educational courses in which learners participate in interactions, and the classification results are as follows.

1) Active learners. Learners in this group tend to learn more Civics courses, their learning ability value is always at a high level, not only learning participation is high, but also learning ability is stronger.

2) Potential learners. The number of Civics courses that learners in this group are involved in learning is significantly less than that of active learners. The learners in this group are still high and almost equal to the active learners. Although they have only studied a small number of Civic and Political Education courses, they are very capable of learning and have a great potential to be tapped.

3) Inactive learners. Learners in this group are not active in the whole learning process, their learning effect is not good, they seldom participate in the learning of the Civic and Political Education courses, and their own ability is weak.

8.1. Adaptation Performance Analysis of Civic Education Content

The performance of active learners, potential learners and inactive learners in accuracy and recall is specifically shown in Fig. 3, Figs. (a) and (b) correspond to accuracy and recall, respectively. From the figure, it can be seen that the accuracy rate (precision@N, N i.e., the number of recommended courses) decreases with the increase of N, in which the accuracy rate of active learners is always the highest, followed by potential learners and inactive learners. Since active learners have richer records of learning behaviors, the Civic Education Content Intelligent Generation and Adaptation System proposed in this paper can more accurately describe their learning status. The recall@N is generally lower for active learners and higher for inactive learners, which is consistent with the characteristics of the three learner groups.

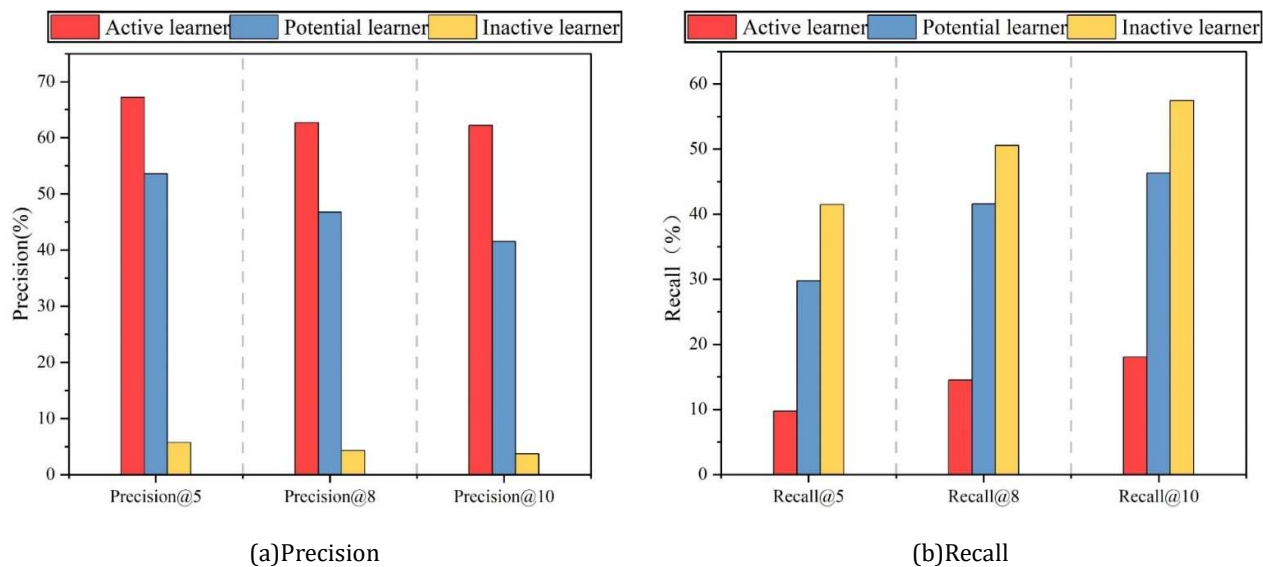


Fig. 3. Precision and recall

Because the change rule of F1 value is complicated, this paper takes more F1 values when $N=10, 20, 30$, and the results are shown in Fig. 4. The $F1@N$ of active learners is increasing, and the larger the value of N is, the more courses the learners watch and the better their performance is. The $F1@10$ and $F1@8$ of potential learners are higher because most of the potential learners watch 1~10 educational courses. As the number of recommended courses increases until it exceeds their ability to accept (e.g., $N \geq 10$), the F1 value starts to gradually decrease from $F1@10$. Inactive learners have

relatively weak ability due to the low number of course viewings (0~4), and the F1@N has been decreasing. From the above analysis, it can be seen that when the length of recommendation list is fixed, the precision@N can be more objective and reasonable to evaluate the accuracy of recommendation algorithm, because this index refers to the proportion of successfully recommended courses in the recommendation list, which is not affected by the identity label of learners.

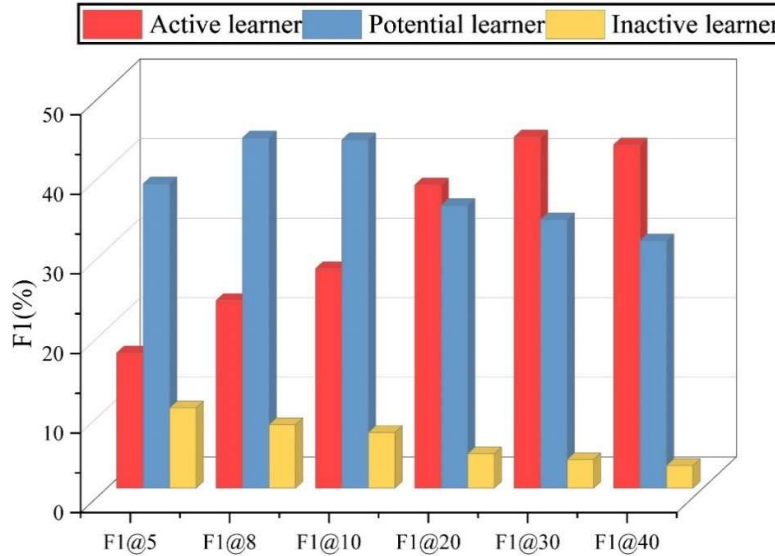


Fig. 4. F1

The cumulative hit rate (hit_ratio@N), i.e., the ratio of the number of successful referrals to the total number of referrals) of the three different learner groups was used to reflect the probability of success of the experiment, as shown in Table 3. It can be seen that the cumulative hit ratio (hit_ratio@N) increases with the increase of N. The probability of successful recommendation is also shown in Table 3, which reflects the probability of success of the experiment. For both active and potential learners, the cumulative hit rate (hit_ratio@N) is very high, where the hit_ratio@5, hit_ratio@8, and hit_ratio@10 of potential learners are higher than those of active learners by 0.74%, 2.79%, and 1.15%, respectively. The reason is that these two types of learners provide enough behavioral logs that can be used to mine their learning rules, which greatly increases the likelihood of recommendation success. On the contrary, the inactive learners' participation and ability are much lower than those of the other two groups, so it is difficult to summarize their learning patterns and accurately recommend educational courses for them.

Table 3. Hit_ratio

Learner group	Hit_ratio@N		
	Hit_ratio@5(%)	Hit_ratio@8(%)	Hit_ratio@10(%)
Active learner	94.47	95.87	97.54
Potential learner	95.21	98.66	98.69
Inactive learner	28.48	32.26	36.01

8.2. Performance Analysis of Personalized Recommendations for Civics Courses

Next, we will evaluate the performance of the intelligent generation and adaptation system of the Civics education content in this paper for personalized recommendation of Civics courses in three groups of learners, and test whether the recommended learning content meets the learners' abilities by examining the adaptivity indicator (adaptivity@N), which refers to the difference between the

average difficulty of the educational courses in the recommended list and the average difficulty of the courses that the learners have already taken). The results of adaptivity@N) are shown in Table 4. For active learners, the average difficulty of the recommended educational courses is 0.08~0.15 higher than the average difficulty of the courses they have actually taken, and for potential learners, the difficulty of the recommended courses is also slightly higher than that of the courses they have actually taken, but the difference between the higher levels ranges from about 0.06 to 0.085, and the difficulty difference is relatively low. For inactive learners, the overall difficulty difference is between -0.01 and 0.015, which is consistent with their weaker ability. Overall, the Civics courses recommended by the system in this paper are in line with the learners' abilities.

Table 4. Adaptivity

Learner group	Adaptivity@5	Adaptivity@8	Adaptivity@10
Active learner	0.145	0.1119	0.0838
Potential learner	0.0823	0.0833	0.0687
Inactive learner	-0.019	-0.017	0.0146

In this section, we will show the difficulty difference between the Civics education courses actually studied by a potential learner in a certain learning time period and the education courses recommended by the system of this paper for him/her, as shown in Fig. 5. It can be seen that the difficulty of the educational courses recommended by the system in this paper is mostly higher than or equal to the difficulty of the actual learning courses, which controls the difficulty of the learning content well within the learner's ability and provides some challenging contents.

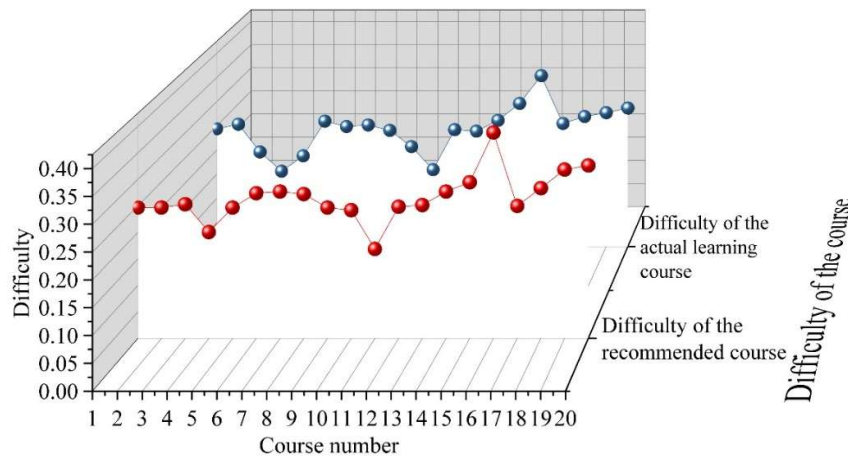


Fig. 5. Course difficulty

9. Conclusion

Based on artificial intelligence content generation technology, this paper constructs its resource automation generation model, civic education knowledge tracking model, civic education course personalized recommendation model in accordance with the intelligent upgrading steps of civic education, and builds the civic education content intelligent generation and adaptation system by combining the model with the model planning system functional modules. The AUC and R2 indicators of this paper's knowledge tracking model in the civic education knowledge tracking experiment are better than those of the AFM model, the standard BKT model, and the extended D-BKT model and T-BKT model name of the BKT model, and it can better simulate the learner's response. The larger the value of N in F1@N for active learners, the more lessons the learners watched. The cumulative hit rate (hit_ratio@N) is used to reflect the success rate of personalized recommended courses in this

experiment The cumulative hit rate of active learners and potential learners are very high, while the cumulative hit rate of inactive learners is less than 40%, which indicates that the lack of behavioral logs of inactive learners poses a difficulty for personalized recommendation of courses. In the results of adaptivity (adaptivity@N), the average difficulty of the recommended educational courses for active learners is 0.08 to 0.15 higher than the average difficulty of the courses they actually take, and about 0.06 to 0.085 higher for potential learners, while the difference for inactive learners ranges from -0.01 to 0.015, and the recommended Civics courses basically match the learners' learning ability.

References

- [1] Oksanen, A., Cvetkovic, A., Akin, N., Latikka, R., Bergdahl, J., Chen, Y., & Savela, N. (2023). Artificial intelligence in fine arts: A systematic review of empirical research. *Computers in Human Behavior: Artificial Humans*, 100004.
- [2] Bolanos, F., Salatino, A., Osborne, F., & Motta, E. (2024). Artificial intelligence for literature reviews: Opportunities and challenges. *Artificial Intelligence Review*, 57(10), 259.
- [3] Borges, A. F., Laurindo, F. J., Spínola, M. M., Gonçalves, R. F., & Mattos, C. A. (2021). The strategic use of artificial intelligence in the digital era: Systematic literature review and future research directions. *International journal of information management*, 57, 102225.
- [4] Doshi, A. R., Bell, J. J., Mirzayev, E., & Vanneste, B. S. (2024). Generative artificial intelligence and evaluating strategic decisions. *Strategic Management Journal*.
- [5] Cacciamani, G. E., Siemens, D. R., & Gill, I. (2023). Generative artificial intelligence in health care. *Journal of Urology*, 210(5), 723-725.
- [6] Corchado, J. M., López, S., Garcia, R., & Chamoso, P. (2023). Generative artificial intelligence: fundamentals. *ADCAIJ: advances in distributed computing and artificial intelligence journal*, 12(1), e31704-e31704.
- [7] Mesko, B. (2023). The ChatGPT (generative artificial intelligence) revolution has made artificial intelligence approachable for medical professionals. *Journal of medical Internet research*, 25, e48392.
- [8] Farrelly, T., & Baker, N. (2023). Generative artificial intelligence: Implications and considerations for higher education practice. *Education Sciences*, 13(11), 1109.
- [9] Jovanovic, M., & Campbell, M. (2022). Generative artificial intelligence: Trends and prospects. *Computer*, 55(10), 107-112.
- [10] Boscardin, C. K., Gin, B., Golde, P. B., & Hauer, K. E. (2024). ChatGPT and generative artificial intelligence for medical education: potential impact and opportunity. *Academic Medicine*, 99(1), 22-27.
- [11] Yu, H., & Guo, Y. (2023, June). Generative artificial intelligence empowers educational reform: current status, issues, and prospects. In *Frontiers in Education* (Vol. 8, p. 1183162). Frontiers Media SA.
- [12] Bosio, E., & Waghid, Y. (2023). Cultivating students' critical consciousness through global citizenship education: Six pedagogical priorities. *Prospects*, 1-12.
- [13] Yu, Y. (2023). Cultivation of Students' Innovative Ability in the Teaching of Ideological and Political Courses in Colleges and Universities. *Curriculum and Teaching Methodology*, 6(12), 45-52.
- [14] Akramova, G. R. (2017). Modern approaches to development critical thinking of students. *Eastern European Scientific Journal*, (5).
- [15] Fenghua, W., & Noordin, Z. M. (2024). The Cultivation of College Students' Sense of Social Responsibility: An Analysis of Student's Involvement in Social Practice. *International Journal of Academic Research in Progressive Education and Development*, 13(1).

- [16] Satria, I. (2022). Cultivating Morals in the Formation of Behavior Student at MTSN 1 Bengkulu City. *MADANIA: JURNAL KAJIAN KEISLAMAN*, 26(2), 237-246.
- [17] Hu, J., & Zhou, X. (2022). Exploring the path of cultivating college students' personal morality in ideological and political theory courses. *Journal of Contemporary Educational Research*, 6(1), 83-88.
- [18] Fonseca, I., Bernate, J., Betancourt, M., Barón, B., & Cobo, J. (2019, October). Developing social responsibility in university students. *In Proceedings of the 11th International Conference on Education Technology and Computers* (pp. 215-218).
- [19] Xia, H., & You, Y. (2023). China: Fostering Students with All-round Attainments in Moral, Intellectual, Physical and Aesthetic Grounding. *In Key Competences and New Literacies: From Slogans to School Reality* (pp. 101-125). Cham: Springer International Publishing.
- [20] Yu, J. (2022). Education for all-round student development: An evaluation of the Compulsory Education Curriculum Program and Standards 2022. *Science Insights*, 41(4), 655-661.
- [21] DEY, T., & GHILDYAL, P. (2021). All-round Development and Comprehensive Assessment. *Journal of Indian Education*, 47(3), 13.
- [22] Miao, M. (2017, December). The All-round Development of Human Being: the Ultimate Value Point of Ideological and Political Education. *In 3rd International Conference on Economics, Management, Law and Education* (EMLE 2017) (pp. 745-751). Atlantis Press.
- [23] Yanan Liu. (2024). Research on Innovative Practices in Teaching Drama and Film Studies under the AIGC. *Journal of Educational Research and Policies*(11),79-82.
- [24] XinLi,XiaoYue,WenYe,YuFu & JiankuiChen. (2024). P-16.2: A Method for Detection of Ink Droplet Volume Distribution Based on Bayesian Theorem. *SID Symposium Digest of Technical Papers*(S1),1536-1540.
- [25] Šarić Grgić Ines,Grubišić Ani & Gašpar Angelina. (2024). Correction: Twenty-Five Years of Bayesian knowledge tracing: a systematic review. *User Modeling and User-Adapted Interaction*(5),1-1.