

Innovation of Virtual Simulation Teaching Strategy Based on Deep Learning Algorithm in Mechanical Manufacturing Experimental Teaching

Jingjing Nie¹, 

¹ School of Mechanical and Electrical Engineering, Wuhan Business University, Wuhan, Hubei, 430000, China

ABSTRACT

Traditional mechanical manufacturing experimental teaching is limited to one teacher demonstrating operations to several students at the same time, which is difficult to take into account and evaluate the differences in knowledge mastery of different students. In order to improve the above teaching defects, firstly, the teaching evaluation of students' experimental level is carried out based on their experimental operation behaviors through K-means clustering. On this basis, a deep learning-based knowledge tracking SAFFKT model is designed to empower and update students' knowledge status. A personalized teaching recommendation method for virtual simulation is proposed based on students' knowledge state, and the hidden semantic matrix decomposition recommendation algorithm for teaching recommendation is improved and implemented. The AUC and ACC of SAFFKT model are significantly higher than that of the comparison model ($p < 0.01$), and it is robust. The F1 value of the recommended experiments was 0.775, indicating a better recommendation effect. The teaching evaluation model achieves accurate classification of students' experimental behavior and yields different learning characteristics of three types of students. Therefore, the innovative work of virtual simulation teaching strategy in this paper is of practical significance.

Keywords: Saffkt; k-means, deep learning, mechanical manufacturing experiment,, virtual simulation teaching

1. Introduction

In order to cultivate innovative and high-quality mechanical manufacturing talents with solid professional theoretical foundation, good practical ability, and adapt to the needs of national economic and social development, the construction of mechanical manufacturing experimental teaching is very necessary [1-2]. However, the traditional high-end machinery to carry out physical experiments

 Corresponding author.

E-mail address: jingjinglinda0324@163.com (J. Nie).

Received 15 January 2024; Revised 25 March 2024; Accepted 30 December 2024; Published Online 16 April 2025.

DOI: [10.61091/jcmcc127b-235](https://doi.org/10.61091/jcmcc127b-235)

© 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/>).

difficult, large investment, especially large-scale mechanical manufacturing process and dangerous type of field experiments are difficult to carry out [3-4]. Simulation of large-scale equipment manufacturing site equipment manufacturing virtual simulation experiment is one of the necessary means of mechanical undergraduate experiments [5].

With the development of computer technology, hardware and software processing capabilities continue to improve, virtual simulation technology to construct the virtual environment is more and more realistic, the sense of experience is more and more real [6-7]. The application of computers in the fields of education and teaching and engineering training is also more and more extensive, and the application level is higher and higher [8]. The Ministry of Education requirements based on the construction of virtual simulation laboratory this task has been proposed to the universities and some universities have been built [9]. At present, the contradiction between the rapid enrichment of teaching content and the relative compression of course credit hours, leading to how to let students do more experimental content under the limited credit hours, higher learning initiative, more solid knowledge mastery needs [10-11]. Therefore, the construction of virtual simulation experimental teaching content suitable for Internet networking, the realization of the integration of the classroom, dormitory and laboratory sharing applications, adapted to the students interested in and at any time and anywhere use of the mobile terminal interactive operation mode, so that the students can conveniently preview, operation, communication, is an effective way to meet this demand [12-15].

The mechanical courses in many Chinese universities have developed many virtual simulation experimental programs based on the in-class experimental contents of many courses, respectively relying on different simulation software [16-17]. For the advanced manufacturing simulation experimental system in the mechanical system, most of the current institutions teachers stay in the demonstration of the cognitive level stage, only play a role in improving the level of perceptual knowledge cognition. At the same time for some of the provision of parameterizable input, process-controllable machining virtual simulation platform, mostly based on the accurate modeling of finite element models and the encapsulation of finite element software to achieve [18-20]. Although appropriate validation virtual experiment simulation can be carried out, the simulation results depend on the accuracy of the finite element model, and at the same time, the numerical computation of the finite element model does not have real-time interactivity, which is unable to meet the requirements of large-scale experiments [21-23].

For the teaching of mechanical manufacturing experiments, an evaluation model of students' proficiency in experimental operations is proposed to determine students' mastery of mechanical manufacturing experiments by quantifying a variety of students' experimental behaviors, which are clustered by K-means clustering. Given that the strategic innovation of virtual simulation teaching cannot be separated from the mastery of students' knowledge state, a deep learning-based knowledge tracking SAFFKT model is designed and optimized. The model consists of a self-attention layer, a forgetting update layer, and a memory reading layer, which respectively realize the functions of knowledge point empowerment, knowledge memory state update, and updating knowledge point mastery matrix. The virtual simulation teaching strategy of personalized recommendation is proposed in accordance with students' knowledge status, and the distribution of knowledge points is portrayed through the knowledge map, knowledge centrality and time effect function to jointly construct the knowledge point association matrix. The time effect function is introduced to propose a hidden semantic matrix decomposition recommendation algorithm suitable for teaching recommendation.

2. Teaching Evaluation Methods for Mechanical Manufacturing Experiments

In the process of the experiment, the students' experimental operation behavior and experimental results there is an inevitable cause and effect relationship, the experimental behavior will affect the experimental results, and the experimental results will also reflect the students' mastery of the experiment, if the students in the end of doing the experiment to come up with a good experimental

results, we can judge that the students have been very skillful in this experiment. Then how to evaluate students' proficiency in mechanical engineering experimental operation needs further research. We first collect sample data that can reflect students' experimental behavior and cluster them on this basis.

2.1. Description of the problem

In the process of the experiment, the students' experimental operation behavior and experimental results there is an inevitable cause and effect relationship, the experimental behavior will affect the experimental results, and the experimental results will also reflect the students' mastery of the experiment, if the students in the end of doing the experiment to come up with a good experimental results, we can judge that the students have been very skillful in this experiment. Then how to evaluate students' proficiency in mechanical engineering experimental operation needs further research. We first collect sample data that can reflect students' experimental behavior and cluster them on this basis.

1) Student sample set. Use $X = \{X_1, X_2, X_3, \dots, X_n\}$ to denote the set of student samples where each object has attributes of m dimensions.

2) Students' experimental behaviors. Determining a student's proficiency in a mechanical engineering experiment is analyzed based on the student's multiple experimental manipulation behaviors. Students' experimental behaviors include students' time for choosing experimental equipment, number of experimental operation steps, number of student experiments, accuracy of experimental equipment placement, accuracy of operation steps, and time for completing experiments, etc., and these six experimental operation behaviors are also the six attributes of students. Numericalizing these behaviors as input data, the proficiency of students is judged through analysis and quantification.

The eigenvalue of the behavioral attribute of each student is denoted as X_{it} ($1 \leq i \leq n, 1 \leq t \leq m$), X_{it} denoting the t th attribute of the i th object.

3) Student experimental proficiency division. The set of student behavioral attribute values is divided into $j = (1 \leq j \leq k)$ classes according to the requirements, and each class contains students with similar experimental proficiency. For K-means algorithm, k clustering centers $\{C_1, C_2, C_3, \dots, C_k\}, 1 \leq k \leq n$ are initialized first [24].

4) Objective function. By calculating the Euclidean distance of each object to the clustering center close to it, X_i denotes the i th object, C_j denotes the j th clustering center, and C_{jt} denotes the t th attribute of the j th clustering center. This is shown in the following equation:

$$dis(X_i, C_j) = \sqrt{\sum_{t=1}^m (X_{it} - C_{jt})^2} \quad (1)$$

2.2. Model structure of students' proficiency in laboratory operations

Based on the above basics, a framework for evaluating students' experimental proficiency was constructed, as shown in Figure 1. Firstly, the required information of students' experimental behaviors was obtained through the experimental environment of the mobile terminal, and then this behavioral information was transformed into computer-recognizable data, and then the preprocessed data set was clustered with K-means. After conducting multiple data training and multiple clustering, it was found that taking the K value of the clusters as 3 was the most appropriate, so the students' experimental proficiency was categorized into high, medium and low. Finally, the student's one time behavioral data information was entered to determine the student's experimental proficiency.

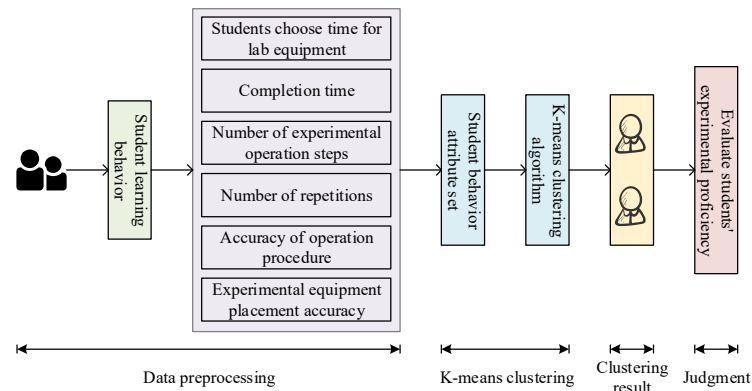


Fig. 1. Model structure of students' experimental operation proficiency

3. Deep learning-based knowledge tracking model design and optimization

The innovation of virtual simulation teaching strategies firstly requires teachers to have a comprehensive knowledge of students' learning status. Dynamic key-value pair memory networks help the knowledge tracking task get a breakthrough achievement in fitting students' learning behavior modeling, which not only improves the computation rate and correctness of knowledge tracking, but also greatly increases its interpretability for the user's behavior model [25]. In order to better fit students' individual characteristics as well as historical practice and revision states, this paper proposes SAFFKT, a knowledge tracking algorithm based on the self-attention mechanism and forgetting memory fitting.

The model includes three key components: an optimized self-attention layer, a forgetting update layer, and a memory reading layer. On the basis of feature fusion, the self-attention layer increases the weight value of the knowledge points in the test questions, which enhances the difference of the knowledge points in the test questions and the proportion of strengths and weaknesses according to the number of different knowledge points and the proportion of knowledge points in the test questions. And two important dynamic update vectors are added to the memory matrix, the forgotten state and the memorized state of the knowledge points. The state quantities will dynamically update the students' knowledge point mastery vectors in the process of learning, and constantly fit the dynamic matrix of students' knowledge point mastery. The role of the forgotten update layer is to obtain the historical sequence data of the student during the learning process and write it into the forgotten knowledge point mastery matrix. After the student's answer, the student's memorized knowledge point mastery matrix is updated based on the memory state quantity in the memory reading layer, whose model is shown in Fig. 2.

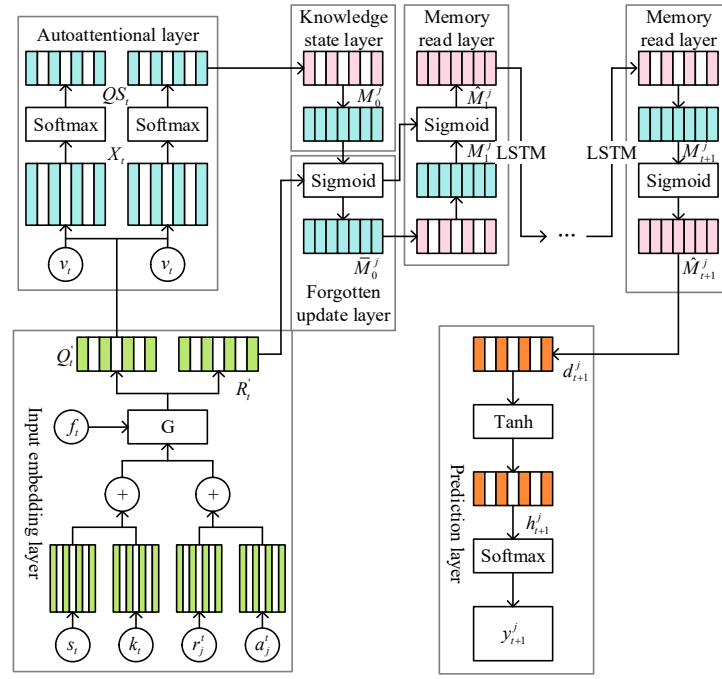


Fig. 2. SAFFKT model structure

3.1. Self-attention mechanism layer optimization

The self-attention layer of SAFFKT is based on feature fusion with the addition of test-knowledge point weight values, which can be interpreted as a constant association and evaluation of knowledge point weight values related to the question based on the student's answer to test question q_t in the process of answering the question. The model will first multiply the test question v_t with the test question-knowledge point embedding matrix X_t in the process of answering the question to obtain the knowledge point embedding vector k_t , and then make it pass through the knowledge weight vector QS_t to obtain the test question-knowledge point weight value w_t , whose formula is as follows:

$$w_t(i) = \text{Softmax}(k_t^T \cdot QS_t(i)) \quad (2)$$

Where, w_t denotes the test-knowledge point weight value, k_t^T denotes the transpose of the knowledge point embedding vector of test v_t , $QS_t(i)$ denotes the knowledge weight vector, and line i of $QS_t(i)$ denotes the i th knowledge point c_i for which the *Softmax* value of the inner product operation between the vectors is taken to compute the attention $w_t(i)$ value, i.e., the weight of the knowledge point c_i in the test v_t . The inner product operation *Softmax* between vectors is calculated as follows:

$$\text{Softmax}(x_i) = \frac{e^{x_i}}{\sum_{j=1}^i x_j} \quad (3)$$

3.2. Oblivion update layer

The process of forgetting update is to help the knowledge tracking task to process the forgetting before the students answer the questions by reading the students' knowledge point mastery matrix M_t^j using the test-knowledge point weights w_t in the self-attention mechanism model mentioned above and adding the forgetting factor in the process of reading. The way is that firstly, according to the student's knowledge point mastery matrix M_t^j at the moment of t , the sigmoid activation function is used to calculate the student's knowledge point forgetting state quantity F_t^k at this time, and then

according to the forgetting state quantity and the test-question-knowledge point weight value, the student's forgotten knowledge point mastery matrix $\bar{M}_t^j$ is calculated, and the formula is as follows:

$$F_t^j = \text{Sigmoid}(E_1 M_t^j + b_1) \quad (4)$$

$$\bar{M}_t^j = M_t^j [1 - w_t F_t^k] \quad (5)$$

Where, F_t^j denotes the amount of forgotten knowledge point state of student j at moment t , $\bar{M}_t^j$ denotes the forgotten knowledge point mastery matrix of student at moment t , M_t^j denotes the knowledge point mastery matrix of student, w_t denotes the test-knowledge point weight value, $\text{Sigmoid}(\cdot)$ denotes the activation function, and E_1 and b_1 denote the weight and bias of the activation function.

3.3. Memory reading layer

The memory readout layer is used to further update the knowledge state based on the student's question answering after the student has finished answering the question. The model will obtain the knowledge increment T_t based on the student's question answering status this time by multiplying the output of the knowledge point state layer with the test-knowledge point embedding matrix X_1 . To further characterize the student's state, the knowledge point weight vector QS_t needs to be spliced with the knowledge increment to obtain the optimized knowledge increment T_t^j , which is given by the following equation:

$$T_t^j = (T_t \oplus QS_t) \cdot M_{t+1}^j \quad (6)$$

Where T_t^j denotes the optimized knowledge increment of student j at moment t , T_t denotes the knowledge increment of test question v_t , QS_t denotes the knowledge weight vector between test question v_t and the covered knowledge points, and M_{t+1}^j denotes the knowledge point mastery matrix of student j at moment $t+1$.

Since the erasure vectors in the forgetting update layer are erased according to the amount of students' knowledge point forgetting state, if the amount of students' knowledge point forgetting state is the same, their erasure vectors are still the same, which will lead to the model forgetting too much content in the process of training and losing the authenticity of the students in the learning process. In order to better fit the forgetting factors in the learning process of students, the model is improved on the time series model by linearly combining the current amount of forgotten state of students' knowledge points F_t with the amount of memorized state R_t^j to obtain the vector of memory updates of knowledge points, which is given by the following formula:

$$R_t^j = w_t(i) \text{sigmoid}(E_2 T_t^j + b_2) \quad (7)$$

$$N_t^j = \omega_1 R_t^j + \omega_2 F_t^j \quad (8)$$

Where, R_t^j denotes the amount of memorized knowledge point state of student j at moment t , w_t denotes the value of test-knowledge point weight, $\text{sigmoid}(\cdot)$ denotes the activation function, T_t^j denotes the optimized knowledge increment of student j at moment t , N_t^j denotes the amount of memorized knowledge point update of student j at moment t , ω_1 & ω_2 denote the value of vector weight, and E_2 & b_2 denote the weight and bias of activation function.

Finally, the student's memorized knowledge point mastery matrix $\hat{M}_{t+1}^j$ is calculated with the following equation:

$$\hat{M}_{t+1}^j = M_{t+1}^j \cdot N_t^j \quad (9)$$

Where, $\hat{M}_{t+1}^j$ denotes student j 's memorized knowledge point mastery matrix M_{t+1}^j at moment $t+1$, denotes student j 's knowledge point mastery matrix at moment $t+1$, and N_t^j denotes

student j 's memorized knowledge point update at moment t .

4. Virtual simulation teaching strategy based on personalized recommendation

After realizing knowledge tracking, a good teaching strategy also needs to arrange appropriate learning or practice contents for students according to their knowledge status. In order to improve the degree of personalization of virtual simulation teaching, this paper takes into account the relationship between each knowledge point and the interaction between users to construct a teaching knowledge point association matrix, and combines this matrix with the use of the law of forgetting knowledge points to improve the traditional hidden semantic matrix decomposition recommendation method.

4.1. Construction of a knowledge point association matrix incorporating multiple factors

The assessment of users' mastery of knowledge points in virtual simulation teaching is often very complex, and it is not possible to analyze the score of the online assessment as the users' mastery of knowledge points in a single way. In this section, a multifactorial knowledge association matrix is constructed by integrating the knowledge map, knowledge centrality and time effect function, and the user's behavioral data in the virtual simulation teaching and the distribution of knowledge points in the assessment test are analyzed in detail, so as to construct a multifactorial knowledge association matrix to accurately assess the user's mastery of knowledge points.

4.1.1. Knowledge maps. In order to assess the distribution and association of knowledge points in virtual teaching tests, this paper introduces the concept of knowledge map. Taking the virtual teaching of mechanical manufacturing as an example, we analyze the relationship between each operation in the simulation experiment step of mechanical manufacturing and the knowledge units in the specific course area, design and construct the knowledge map to react to the user in the form of visualization of the knowledge linkage, so that the user can understand the distribution of each knowledge point and the knowledge linkage according to the map, and find the resources that meet their needs [26].

4.1.2. Degree of centrality of knowledge. In this paper, the concept of centrality is introduced into the field of online teaching resources recommendation, which is used to describe the degree of association between each knowledge node and the importance of each knowledge node. In this paper, we consider that knowledge nodes with large centrality have more connections with the surrounding knowledge nodes, on the contrary, for knowledge points with small centrality, which are relatively less connected with other knowledge points, the importance of such knowledge points is relatively much lower. So similar to the definition of network centrality, this paper defines the centrality Center of a knowledge point is expressed as:

$$Center_{Degree}(K_i) = \frac{\sum_j relation_{ij}}{sum - 1}, i \neq j \quad (10)$$

where $relation_{ij}$ is the number of associations between knowledge point i and knowledge point j , and sum is the sum of the number of knowledge nodes in this knowledge map.

4.1.3. Knowledge point correlation matrix incorporating multiple factors. In this paper, we establish a fusion multifactor-based correlation matrix of users' knowledge point mastery. Assuming that S user $U = \{u_1, u_2, \dots, u_s\}$ participated in the same virtual teaching test, which required a total of M test steps $S = \{s_1, s_2, \dots, s_m\}$ to be completed, in which these test steps examined a total of N knowledge points $I = \{I_1, I_2, \dots, I_N\}$, and that user u_{helper} assisted user u_{test} to complete the experimental step $s_{difficult}$ during the experiment, the relevant factors for constructing the correlation matrix of user knowledge point mastery include:

- 1) The rate of missing points. The miss rate describes the miss score of user u on step s of the experiment. If it passes completely, record the lost score rate for that step $g_{us} = 0$, otherwise $g_{us} = 1$. The aggregation yields the distribution matrix $G_{s \times m}$ of the loss of score for user u on step s .
- 2) Assisted Participation Impact Factor. If user u_{helper} assists user u_{test} in completing step $s_{difficult}$ of the experiment, the score of the assisted test user u_{test} for the knowledge points contained in step $s_{difficult}$ will be improved, and the degree to which the score is affected by the assistance is defined here γ as the assistance participation influence factor. So the distribution matrix of the scores lost by the users in the test steps $G_{s \times m}^*$ is obtained by adding the assistance participation factor into consideration:

$$g_{u,s}^* = 1 - (1 - g_{u,s}) \times \gamma_{u,s} \quad (11)$$

$$\gamma = 1 - e^{-h} \quad (12)$$

Where $g_{u,s}$ denotes the loss of points by user u at step s , $g_{u,s}^*$ denotes the loss of points by user u at step s with assistance from other users, γ is the assistance participation influence factor, and h is the number of times assisted.

- 3) Step-related knowledge. In virtual simulation teaching tests, there is often more than one knowledge point associated with each test step, so here, experimental step S is disassembled into multiple knowledge points, and it is stipulated that if knowledge point I is examined in step S , it will be recorded as 1, otherwise it will be recorded as 0. This results in the distribution matrix of knowledge points in the experimental steps $K_{M \times N}$.

- 4) Importance of knowledge points. The matrix $K_{M \times N}$ is improved and optimized using the knowledge point centrality formula and the original matrix is repopulated after normalization to obtain the matrix $K_{M \times N}^*$ of the distribution of the importance of knowledge points in each experimental step.

Based on the above correlation factors, the product of the failure rate of an experimental step and the knowledge points included in that step is used to obtain a weighted matrix R_{ai} of the failure of user u on knowledge point i :

$$R_{ui} = G_{us}^* \times K_{si} \quad (13)$$

The larger the value of r_{in} , the more points user u loses in knowledge point i , i.e., the lower the user's mastery of the knowledge point, and vice versa, the better the mastery.

4.2. Matrix Decomposition Recommendation Algorithm Based on Improved Hidden Semantic Modeling

In this section, the traditional matrix decomposition based on the cryptosemantic model is improved by combining the knowledge forgetting scale, and an improved cryptosemantic matrix decomposition method based on the knowledge point association matrix is proposed as an input to the algorithm.

- 4.2.1. Time effect function. In this paper, a time effect function is introduced to assess the trend of users' mastery of knowledge points over time. However, after the direct introduction of the forgetting curve, the time effect of the two time points closer to the virtual teaching test interval will be over-expanded, while the time effect of the two time points farther away from the test interval will be ignored because of its small size. Based on the above reality and combined with the Ebbinghaus forgetting law and the user's mastery of knowledge, we define the time effect function as follows:

$$f(t) = \mu + (1 - \mu) \times (1 - e^{-(t-t_0)}) \quad (14)$$

where t indicates the current time, t_0 indicates the time when the user takes the virtual teaching test. μ is the time effect influence factor, indicating the degree of influence of the time effect on the user, used to better match the user's knowledge point mastery, and $\mu \in [0, 1]$, if $\mu = 0$, it is considered

that the user's knowledge mastery of the forgetting process completely follow the time effect function, and vice versa, if $\mu=1$, it does not comply. In order to simplify the arithmetic process, this paper defines the time effect function is a monotonically increasing function model, with the passage of time the user forgets the content of the knowledge points will be increased, the user's loss of points in the specific knowledge points will gradually increase, reflecting the user for the degree of mastery of the knowledge points will show a trend of continuous decay.

4.2.2. Recommendation algorithms incorporating time effects. Based on the implicit semantic model [27], Definition:

$$P = U_{m \times r} \Sigma_{r \times r} \quad (15)$$

$$Q = V_{n \times r} \quad (16)$$

So the user-knowledge mastery matrix R can be expressed as:

$$R = PQ^T \quad (17)$$

Where $P_{m \times f}$ is the user's implicit semantic matrix about the attribute categorization of knowledge points, and $Q_{n \times f}$ is the proportionate weight of each knowledge point in the attribute categorization. So the user's mastery of knowledge points $\hat{r}_{ui} = R(u, i)$ can be transformed into:

$$\hat{r}_{ui} = R(u, i) = \sum_{f=1}^F p_{uf} q_{if} \quad (18)$$

Among them, the hidden class $f \in (1, F]$, $p_{ui} = P(u, i)$, $q_{if} = Q(i, f)$. In this paper, we assume that the difference between the real results of the user's knowledge mastery and the predicted results of the knowledge points obeys a Gaussian distribution, and we utilize the constructed knowledge point correlation matrix integrating multi-factors as the distribution matrix of the knowledge mastery of the existing users to calculate, so as to derive the relevant parameters of the hidden class matrices P and Q . In this paper, the root mean square error formula is used to evaluate the degree of conformity between the prediction results and the user's actual mastery, and the loss function is denoted as $C(p, q)$:

$$RMSE = \sqrt{\frac{\sum_{u,i \in T} (r_{ui} - \hat{r}_{ui})^2}{T}} \quad (19)$$

$$C(p, q) = \sum_{(u,i) \in Train} (r_{ui} - \hat{r}_{ui})^2 = \sum_{(u,i) \in Train} \left(r_{ui} - \sum_{f=1}^F p_{uf} q_{if} \right)^2 \quad (20)$$

In order to prevent the problem of overfitting during the learning process, it is necessary to add an overfitting term to Eq. (20), where λ is the regularization parameter can be determined experimentally, then the equation is transformed into:

$$C(p, q) = \sum_{(u,i) \in Train} (r_{ui} - \sum_{f=1}^F p_{uf} q_{if})^2 + \lambda (\|p_u\|^2 + \|q_i\|^2) \quad (21)$$

Considering the trend of the user's mastery of knowledge points over time, the time effect function (21) is introduced to optimize the user-knowledge point mastery matrix r_w , and the process of the user's forgetfulness of knowledge points over time $f_{wi}(t)$ is taken into account to obtain the loss-of-score matrix r_{wi}^* , which is weighted and corrected using the time effect:

$$r_{ui}^* = r_{ui} \times f_{ui}(t) \quad (22)$$

The loss function is then evaluated using the root mean square error, which can be transformed into:

$$\begin{aligned} C(p, q) &= \sum_{(u,i) \in \text{Train}} f(t)(r_{ui} - \hat{r}_{ui})^2 \\ &= \sum_{(u,i) \in \text{Train}} f(t)(r_{ui} - \sum_{f=1}^F p_{uf} q_{if})^2 \\ &= \sum_{(u,i) \in \text{Train}} f(t)(r_{ui} - \sum_{f=1}^F p_{uf} \times q_{if})^2 + \lambda \|p_u\|^2 + \lambda \|q_i\|^2 \end{aligned} \quad (23)$$

The loss function $C(p, q)$ is optimized using gradient descent (SGD) and taking partial derivatives of p_{uf} and q_{if} yields the following results after several iterations of optimization:

$$\frac{\partial C}{\partial p_{uf}} = -2f(t)(r_{ui} - \sum_{f=1}^F p_{uf} q_{if})q_{if} + 2\lambda p_{uf} \quad (24)$$

$$\frac{\partial C}{\partial q_{if}} = -2f(t)(r_{ui} - \sum_{f=1}^F p_{uf} q_{if})p_{uf} + 2\lambda q_{if} \quad (25)$$

$$p_{uf} = p_{uf} + \alpha((r_{ui} - \sum_{f=1}^F p_{uf} q_{if})q_{if} + \lambda p_{uf}) \quad (26)$$

$$q_{if} = q_{if} + \alpha((r_{ui} - \sum_{f=1}^F p_{uf} q_{if})p_{uf} + \lambda q_{if}) \quad (27)$$

Where, α is the learning efficiency, the larger the value of α means that the gradient descent is also faster, the value of which needs to be obtained through a number of experiments, here $\alpha = 0.02$. Through the above to carry out a number of iterations, you can obtain one of the parameters p_{uf} and q_{if} , and then predict the user's mastery score of the knowledge points.

5. Example analysis

5.1. Application of the experimental operation evaluation model

5.1.1. Experimental environment and data sets. The experimental environment is Windows 10 64-bit Professional operating system, the main processor is Intel® Core™ i7-9750H, the host memory is 16GB, 1TB SSD hard disk. The IDE development environment is PyCharm 2019, the programming language is python3.7, curve fitting is done using Matlab 2018b in the cftool tool, and MySQL 5.7 for data preprocessing.

This experiment uses a virtual experimental dataset, including 1655 records of experimental operation behaviors from 528 learners. The characteristics of experiment operation behavior are mainly analyzed in terms of the number of operations, the correct rate of operation, the difficulty of operation, the length of operation, and the learning situation. The mechanical manufacturing experiment video viewing length and practice performance are shown in Fig. 3(a) and (b).

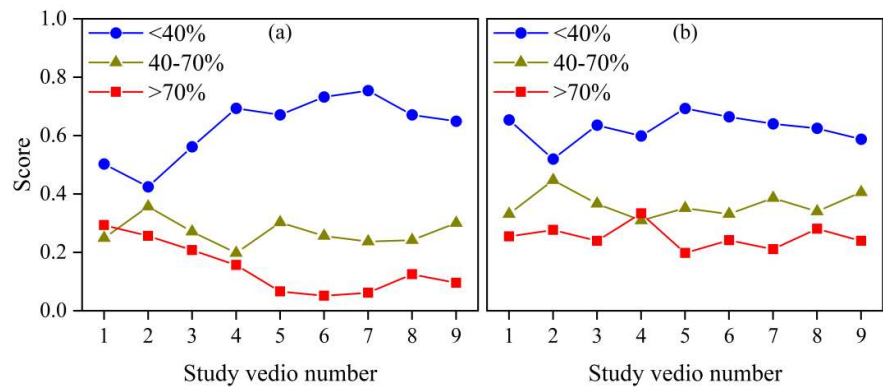


Fig. 3. Part of video watching and practice performance distribution

In the experiment, the five selected attributes are pre-processed, and according to the actual situation, this experiment uses multiple regression analysis to set the corresponding weighting for each attribute feature, and the specific data applied to the experiment are obtained after the secondary calculation. Multiple regression analysis is the use of normalized data to construct a multiple linear regression equation, the significance of the coefficients and the correlation test of the factors, and then select the factors that have the greatest impact on the dependent variable. In this case, the dependent variable is set to be the proficiency of the experimental operation. The final weight values of α for the selected attributes were given as 23%, β as 26%, γ as 17% and δ as 11%. The final pre-processed operation record attribute dataset is obtained as shown in Table 1.

Table 1. Test operation record attribute feature

Sample code	Operation times(OT)	Operation correctness rate	Operation duration(OD)	Learning status(LS)
1	0.9	0.54	0.62	0.72
2	0.8	0.63	0.75	0.82
...				
3	0.65	0.71	0.93	0.83
322	0.58	0.64	0.83	0.91

5.1.2. Analysis of experimental results. The K-means clustering algorithm was applied to the dataset in this experiment, where the accuracy and stability of the algorithm in this experiment was judged based on the number of iterations and the distance between the clustering results and the final clustering center. Table 2 shows the clustering accuracy of the dataset.

From the data in the table, it can be seen that with the increase in the number of iterations, the change in the clustering center becomes smaller and smaller until it finally tends to almost zero. Therefore, the use of K-means algorithm for clustering learner behavior is characterized by a small number of iterations, high stability, and high accuracy of clustering results. After 9 iterations of the clustering algorithm, the change in the value of the objective function tends to stabilize, and at this time, the clustering center no longer changes significantly, and the clustering process is over, so the clustering results after 9 iterations are finally selected as the final clustering results.

Table 2. Data clustering accuracy

Iteration times	Cluster error
2	896.1165764
3	841.7678105
5	821.0896262
7	799.0207356
9	799.5293469
11	799.8121116

The number of clusters expected to be obtained in this experiment is 3. After K-means clustering, the 528 learners are divided into 3 classes based on the experimental operation behaviors, and the clustering results as in Table 3 are obtained.

Among them, the clustering center division of attribute LS is not obvious, which is due to the fact that the learners' viewing of the video and practice of the exercises are basically the same, and the combined scores of the two items do not differ much, which makes the displayed learning situation not obviously differentiated. According to the results of the clustering center table, the learners' experimental operation skills are divided into three levels, which are high, medium and low, corresponding to the three levels of being able to apply flexibly, applying correctly and not being able to apply in the experimental operation, respectively. Then the number of samples belonging to the three clusters was counted. The number of learners who were able to complete the experimental operations flexibly according to their knowledge of mechanical engineering was 125, and learners belonging to this category generally had better overall academic performance, higher rates of correct experimental operations, and shorter and fewer operations. The number of learners who could correctly complete the relevant experimental operations according to their knowledge of mechanical engineering was 162. The number of learners who were not able to complete the experiments correctly according to the theoretical knowledge of mechanical engineering was 241, and the learners belonging to this category were characterized by a low rate of correct experimental operations and a long operation time.

Table 3. Cluster center

Cluster	Cluster center				Sample size in clusters	Operating skill level
	OT	OC	OD	LS		
Cluster 1	0.4	0.85	0.4	0.85	125	High
Cluster 2	0.7	0.5	0.7	0.5	162	Medium
Cluster 3	0.9	0.4	0.9	0.4	241	Low

5.2. Knowledge tracking performance experiments

The effectiveness of SAFFKT is validated on four real education-based public datasets, the effectiveness of the proposed model is verified by comparing the effectiveness of the SAFFKT model with other state-of-the-art knowledge-tracking models on student achievement prediction, the stability and high accuracy of the model's performance on sequences of different lengths are verified by robustness experiments, and the interpretability of the model is verified by visualization experiments.

5.2.1. Baseline data sets. Four publicly available educational datasets were used in this experiment, namely ASSISTment12, ASSISTment17, Slepemapy.cz, and Junyi, and their basic information is shown in Table 4.

Table 4. Dataset Statistic

Dataset	The number of students	The number of knowledge point	The number of exercises	Interaction times
ASSISTment12	24.6k	255	52.8k	2632.6k
ASSISTment17	1.8k	112	3.8k	933.8k
Slepemapy.cz	82.3k	1462	2.5k	9779.4k
Junyi	176.6k	45	711	25658.4k

5.2.2. Experimental setup. Five-fold cross-validation was used in the experiments. For each fold, 80% of the sequences in the dataset were split into a training set (60%) and a validation set (20%), with the remaining 20% serving as a test set. For the comparison algorithms, their original code (DKT, AKT, HawkesKT) or, in the absence of code, a replication based on the original paper (DKT+Forgetting) was used and their parameters were tuned according to the original code or the original paper, respectively. For SAFFKT, hyperparameters were tuned according to the effect of the validation set.

5.2.3. Overall performance of the experiment. Table 5 lists the performance of SAFFKT and all comparison methods, as well as the two extension methods in predicting students' future grades on four datasets and two metrics. The average of the results under 5-fold cross-validation is reported here. In this case, the optimal results on each metric in each dataset are marked using bold, and the suboptimal results in the comparison algorithms are marked using italics and underlining. Meanwhile, the T-test significance test was performed on the optimal and suboptimal results, with the * sign indicating that the optimal results outperform the suboptimal results at a significance level of $p < 0.01$, and the - sign that there is no significant gap between the optimal and suboptimal results. It can be seen from Table 5:

- 1) The SAFFKT method significantly outperformed all baselines on all four datasets for the primary AUC predictor. It also passed the significance test at the $p < 0.01$ level when tested for significance against the suboptimal algorithm, demonstrating that SAFFKT was statistically superior to the comparison algorithm.
- 2) On the secondary ACC predictors, the SAFFKT method significantly outperformed all baselines on

three datasets and slightly outperformed the sub-optimal algorithm on one dataset (Junyi). By performing a T-test significance test with the suboptimal algorithm, SAFFKT passes the significance test compared to the suboptimal algorithm on the three datasets of ASSISTment12, ASSISTment17, and Slepemapy.cz at the level of $p < 0.01$, which proves that the proposed method outperforms the comparison algorithm on all three datasets in a statistically significant way. Moreover, there is no significance gap compared to the sub-optimal results on one dataset, and both of them have comparable performance in a statistically significant way.

Table 5. Overall performance of SAFFKT

Datasets	ASSISTment12		ASSISTment17		Slepemapy.cz		Junyi	
Metrics	ACC	AUC	ACC	AUC	ACC	AUC	ACC	AUC
DKT	0.729	0.691	0.658	0.705	0.795	0.717	0.840	0.703
DKT-V	0.757	0.792	0.691	0.753	0.801	0.756	0.837	0.754
DKT+Forgetting	0.748	0.715	0.684	0.702	0.799	0.767	0.847	0.743
AKT	0.751	0.787	0.702	0.766	0.802	0.779	0.846	0.777
HawkesKT	0.735	0.775	0.668	0.711	0.806	0.739	0.849	0.764
SAFFKT	0.769*	0.815*	0.745*	0.809*	0.812*	0.789*	0.852-	0.784*

5.2.4. Robustness verification. The length of the maximum sequence is an important hyperparameter in the modeling of knowledge tracking models, and the data it corresponds to in practical applications is the number of answers in the students' history. Previous research works do not have a standardized setting for the maximum sequence length, which varies from 50, 100, 200 and other values. In order to verify the robustness of the performance of SAFFKT under different maximum sequence lengths, the variation of the performance of SAFFKT and all baselines under different maximum sequence lengths is reported on the dataset with the largest average sequence length, ASSISTment17 (about 500), and the experimental results are shown in Fig. 4. As can be seen from the figure, SAFFKT exhibits higher stability in AUC and ACC metrics as the sequence length varies.

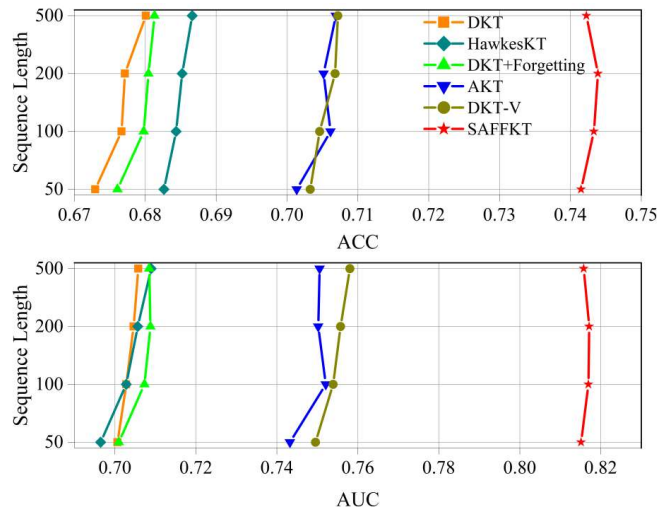


Fig. 4. Robustness analysis based on ASSISTment17

5.2.5. Visualization experiments. In order to verify the interpretability of the proposed method in simulating the student learning process, the change in the knowledge mastery of a student in the test set of the ASSISTment12 dataset is visualized as shown in Figure 5. The figure visualizes the change in the mastery level of a student in the test set for five knowledge points in 30 consecutive steps of learning.

The student's knowledge mastery rises and the color of the subject's heat map tends to be lighter. Otherwise, when the time interval is particularly long, the change in the student's knowledge mastery is mainly due to forgetting, the student's knowledge mastery decreases, and the color of the subject heat map becomes darker.

1) As a result of the learning effect, the knowledge mastery of all observed knowledge points increased significantly on those responses with particularly short time intervals (e.g., from step 0 to step 3 and from step 6 to step 9).

2) The degree of knowledge mastery of all observed knowledge points decreases on those responses with particularly long time intervals (e.g., from step 4 to 5, from step 10 to 11, from step 16 to 18, and from step 26 to 29), which is due to the forgetting effect. Overall, the above analysis clearly verifies the validity and interpretability of the proposed method for modeling students' learning and forgetting processes.

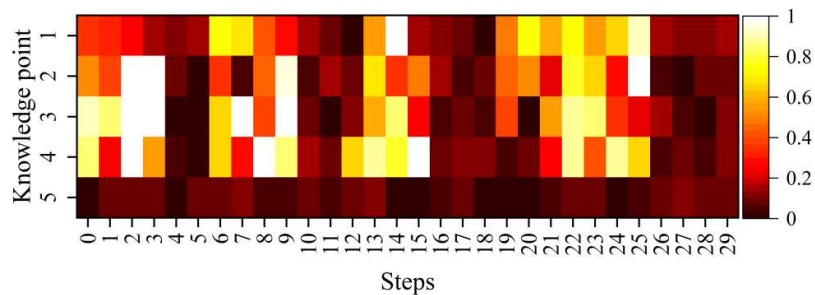


Fig. 5. Visualization of a student's knowledge mastery

5.3. Example analysis of improved recommendation algorithms

In this paper, experiments are conducted on Tensorflow development framework, using python for data preprocessing and algorithm implementation. Accuracy and recall are chosen to evaluate the performance of the recommendation algorithm, in addition, F-Measure method is used for comprehensive evaluation. F-Measure theoretical perspective is the weighted average of accuracy and recall, the higher the F value, the better the recommendation effect. Based on the experimental results, the parameters of the algorithm for data dimensionality reduction using self-encoder in this paper are set as follows: the number of self-encoder layers is 4, and the number of nodes per layer is set to 400.

Take this recommendation algorithm to determine the optimal N value, that is, the number of products in the recommendation list. The relationship between the number of recommendations N and the algorithm accuracy and recall is shown in Figure 6, where the horizontal coordinate is the number of recommendations and the vertical coordinate is the accuracy or recall. The accuracy rate of the algorithm will increase with increasing value of N within a certain range and then decrease. The recall of the algorithm increases with increasing N value. The recommendation is optimal when the point (12) with peak accuracy is selected.

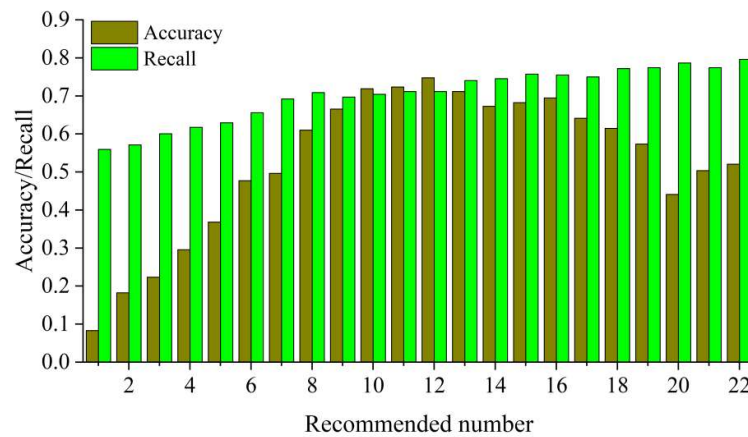


Fig. 6. The relationship between the number of recommendation and accuracy/recall

Figure 7 shows the F1 value of this paper's recommendation algorithm. In this study, 30 users are selected, and by calculation, the average comprehensive index of the users is about 0.775, which means that the comprehensive index of the recommendation result is high, indicating that the recommendation performance of this algorithm is very good.

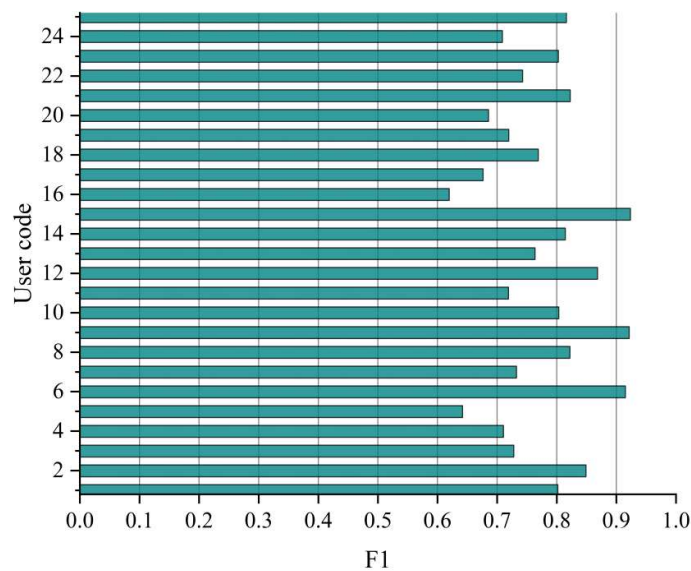


Fig. 7. The F1 value of our recommended algorithm

6. Conclusion

In this paper, after the clustering analysis of the mechanical manufacturing experiment level, based on the deep learning algorithm, the knowledge tracking model and the personalized recommendation method with improved implicit semantic model are designed respectively to innovate the virtual simulation teaching strategy in an all-round way. In this paper, the AUC and ACC predictors of SAFFKT knowledge tracking model are better than other baseline models, and most of the differences show statistical significance ($p < 0.01$), in addition, with the change of sequence length, SAFFKT shows higher stability. Conducting algorithmic recommendation experiments, the average F1 value of 30 users is 0.775, indicating that the method of this paper has a better recommendation effect. The K-means clustering of the teaching evaluation model divides the students into three experimental operation skill levels of “high - medium - low”, with 125, 162 and 241 students in each level respectively. It can be seen that this paper has accomplished the knowledge tracking, personalized teaching and scientific

teaching evaluation of mechanical manufacturing experiments, providing a new strategy for mechanical manufacturing experimental teaching.

Funding

- 1) Hubei Provincial Education Science Planning Project, Research on Generative Teaching Design Based on OBE Concept and Cognitive Load Theory (2021GB074).
- 2) Ministry of Education Industry University Cooperation Collaborative Education Project, Project Name: Construction of Mechanical Manufacturing Training Practice System Based on Domestic Industrial Foundation Core Platform, Project Number: 231005377131038.

References

- [1] Chowdhury, H., Alam, F., & Mustary, I. (2019). Development of an innovative technique for teaching and learning of laboratory experiments for engineering courses. *Energy Procedia*, 160, 806-811.
- [2] Yang, Z., Li, Z., Xiao, J., Tang, H., Du, B., & Wang, L. (2022, December). Research on Practical Teaching System of Intelligent Manufacturing Engineering Major Considering the Improvement of Humanistic Quality. In *2022 International Conference on Educational Innovation and Multimedia Technology (EIMT 2022)* (pp. 331-337). Atlantis Press.
- [3] Wermann, J., Colombo, A. W., Pechmann, A., & Zarte, M. (2019). Using an interdisciplinary demonstration platform for teaching Industry 4.0. *Procedia Manufacturing*, 31, 302-308.
- [4] Mourtzis, D., Zogopoulos, V., & Vlachou, E. (2018). Augmented reality supported product design towards industry 4.0: a teaching factory paradigm. *Procedia manufacturing*, 23, 207-212.
- [5] Chen, W. P., Lin, Y. X., Ren, Z. Y., & Shen, D. (2021). Exploration and practical research on teaching reforms of engineering practice center based on 3I-CDIO-OBE talent-training mode. *Computer applications in engineering education*, 29(1), 114-129.
- [6] Salah, B., Abidi, M. H., Mian, S. H., Krid, M., Alkhalefah, H., & Abdo, A. (2019). Virtual reality-based engineering education to enhance manufacturing sustainability in industry 4.0. *Sustainability*, 11(5), 1477.
- [7] Chong, S., Pan, G. T., Chin, J., Show, P. L., Yang, T. C. K., & Huang, C. M. (2018). Integration of 3D printing and Industry 4.0 into engineering teaching. *Sustainability*, 10(11), 3960.
- [8] Madsen, O., & Møller, C. (2017). The AAU smart production laboratory for teaching and research in emerging digital manufacturing technologies. *Procedia Manufacturing*, 9, 106-112.
- [9] Kapilan, N., Vidhya, P., & Gao, X. Z. (2021). Virtual laboratory: A boon to the mechanical engineering education during covid-19 pandemic. *Higher Education for the Future*, 8(1), 31-46.
- [10] Huang, T. C., & Lin, C. Y. (2017). From 3D modeling to 3D printing: Development of a differentiated spatial ability teaching model. *Telematics and Informatics*, 34(2), 604-613.
- [11] Ford, S., & Minshall, T. (2019). Invited review article: Where and how 3D printing is used in teaching and education. *Additive Manufacturing*, 25, 131-150.
- [12] Zhang, W., Zhang, S., Guo, S., Yang, Y., & Chen, Y. (2017). Concurrent optimal allocation of distributed manufacturing resources using extended teaching-learning-based optimization. *International Journal of Production Research*, 55(3), 718-735.
- [13] Liao, J. Y. (2022). Teaching methods of power mechanical engineering based on artificial intelligence. *Kinetic Mechanical Engineering*, 3(3), 54-61.
- [14] Hernández-de-Menéndez, M., Vallejo Guevara, A., & Morales-Menendez, R. (2019). Virtual reality

- laboratories: a review of experiences. *International Journal on Interactive Design and Manufacturing (IJIDeM)*, 13, 947-966.
- [15] Scaravetti, D., & Doroszewski, D. (2019). Augmented Reality experiment in higher education, for complex system appropriation in mechanical design. *Procedia Cirp*, 84, 197-202.
 - [16] Perini, S., Luglietti, R., Margoudi, M., Oliveira, M., & Taisch, M. (2018). Learning and motivational effects of digital game-based learning (DGBL) for manufacturing education-The Life Cycle Assessment (LCA) game. *Computers in Industry*, 102, 40-49.
 - [17] Li, H., Öchsner, A., & Hall, W. (2019). Application of experiential learning to improve student engagement and experience in a mechanical engineering course. *European Journal of Engineering Education*, 44(3), 283-293.
 - [18] Zhou, J., & Yao, X. (2017). Hybrid teaching-learning-based optimization of correlation-aware service composition in cloud manufacturing. *The International Journal of Advanced Manufacturing Technology*, 91, 3515-3533.
 - [19] Stern, A., Rosenthal, Y., Dresler, N., & Ashkenazi, D. (2019). Additive manufacturing: An education strategy for engineering students. *Additive Manufacturing*, 27, 503-514.
 - [20] Mavrikios, D., Georgoulas, K., & Chrysosouris, G. (2018). The teaching factory paradigm: Developments and outlook. *Procedia Manufacturing*, 23, 1-6.
 - [21] Grodotzki, J., Ortelt, T. R., & Tekkaya, A. E. (2018). Remote and virtual labs for engineering education 4.0: achievements of the ELLI project at the TU Dortmund University. *Procedia manufacturing*, 26, 1349-1360.
 - [22] Chen, J. Y., Tai, K. C., & Chen, G. C. (2017). Application of programmable logic controller to build-up an intelligent industry 4.0 platform. *Procedia Cirp*, 63, 150-155.
 - [23] Ma, J., Jaradat, R., Ashour, O., Hamilton, M., Jones, P., & Dayarathna, V. L. (2019). Efficacy investigation of virtual reality teaching module in manufacturing system design course. *Journal of Mechanical Design*, 141(1), 012002.
 - [24] Chenggang Li, Jiajun Zhang, Haifeng Zhang & Jiarui Wang. (2024). A clustering method for charging behaviors of electric vehicle users based on an improved k-means algorithm. *Journal of Physics: Conference Series*(1), 012029-012029.
 - [25] Sun Xia, Zhao Xu, Li Bo, Ma Yuan, Sutcliffe Richard & Feng Jun. (2021). Dynamic Key-Value Memory Networks With Rich Features for Knowledge Tracing. *IEEE transactions on cybernetics*.
 - [26] Marta Osório de Matos, Elzbieta Bobrowicz Campos, Rosa Silva, Francisca Castanheira & Diana Santos. (2024). Protocol: Complementarity between informal care and formal care to adults: Knowledge mapping through a scoping review of the literature. *Campbell systematic reviews*(4), e70004.
 - [27] Li Ruonan & Xiao Luo. (2023). Latent factor model for multivariate functional data. *Biometrics*(4), 3307-3318.