

Designing a Dynamic Development Mechanism for Young Teachers' Competence in Private Applied Colleges and Universities Using Reinforcement Learning Theory

Hui Luo¹, 

¹ Academic Affairs Office, Geely University, Chengdu, Sichuan, 641423, China

ABSTRACT

The article aims to accelerate the growth and progress of young teachers in private applied colleges and universities and improve their teaching ability, combining with the knowledge graph, and designing a recommended algorithm based on deep reinforcement learning to improve teachers' ability. Firstly, the growth and progress process of young teachers in private applied colleges and universities is defined as a dynamic development process, i.e., for different latitude abilities such as teacher ethics, professional knowledge, preteaching preparation, communication and cooperation, teaching ability training needs to be carried out gradually and in a certain order. Then the Knowledge Graph Teacher Competency Enhancement Recommendation Algorithm (KGDR) based on deep reinforcement learning and knowledge graph algorithm is constructed by combining deep reinforcement learning and knowledge graph algorithm. When performing top- k recommendation, the diversity value of the model at $k = 20$ is 0.7876, and the model can provide more diverse paths for teacher ability improvement. After the application of the dynamic development mechanism of young teachers' competence based on KGDR, the competence improvement of young teachers is significant and can reach the grade of "excellent". The mechanism designed in this paper can be used as a reference for other colleges and universities.

Keywords: deep reinforcement learning; knowledge mapping; dynamic development mechanism; young teachers

1. Introduction

With the popularization and development of higher education, applied private undergraduate colleges and universities are playing an increasingly important role in China's education system. Different from traditional theoretical institutions, applied private undergraduate colleges and universities pay more attention to the cultivation of students' practical ability and employment needs [1-4]. And as the

 Corresponding author.

E-mail address: LuoHui_19820922@163.com (H. Luo).

Received 05 March 2024; Revised 20 April 2024; Accepted 05 December 2024; Published Online 15 April 2025.

DOI: [10.61091/jcmcc127a-280](https://doi.org/10.61091/jcmcc127a-280)

© 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/>).

proportion of young intellectuals in private colleges and universities young teachers in the teaching force is increasing, their quality directly affects the level of higher education [5-6]. Therefore, the establishment of a high-quality young teacher team is not only conducive to the promotion of sustainable development of colleges and universities, but also conducive to the growth of individual teachers [7-9]. And reinforcement learning as a branch of machine learning, its goal is through the interaction between the intelligent body and the environment, so that the intelligent body can learn the optimal strategy from it through trial and error [10-11], it is of great significance to design the dynamic development mechanism of young teachers' competence in private applied colleges and universities by utilizing reinforcement learning.

In the path of teachers' ability development, individuals should focus on the improvement of professional knowledge and teaching ability. Firstly, teachers need to have solid subject knowledge and teaching theory knowledge. Secondly, teachers need to have good teaching skills, including designing teaching content, organizing teaching activities, guiding students' learning, evaluating students' learning, etc. [12-15]. Furthermore, teachers need to continuously improve their research ability in education and teaching, participate in education and teaching reforms, and constantly explore and innovate to promote the development of the teaching profession [16-17]. In addition, in the process of teachers' competence development, there are some key factors that play a key role in teachers' growth and development, namely, perfect teacher training system teachers' personal development planning and orientation [18-20]. Based on their own development planning and orientation, teachers purposefully and actively participate in all kinds of training and learning activities, and continuously learn and improve their professionalism [21-22].

The study introduces the dynamic system theory into the field of young teacher development in private applied colleges and universities, and proposes that the key to the dynamic development of teacher competence is to make teacher competence improvement recommendations. Accordingly, a deep reinforcement learning-based knowledge mapping teacher competency improvement recommendation algorithm is designed. The algorithm combines the knowledge graph of young teachers' competence and attention deep reinforcement learning to carry out teacher competence improvement path recommendation. The recommendation ability of the recommendation algorithm is examined by designing comparative experiments. And with reference to the index system and weights designed in related literature, the designed dynamic development mechanism is empirically analyzed using fuzzy comprehensive evaluation method.

2. Mechanisms for the dynamic development of young teachers' competencies in private applied colleges and universities

As an important part of the faculty of private applied colleges and universities, the dynamic development of young teachers' abilities is of great significance to the improvement of the teaching quality and research level of the school. Private applied colleges and universities take the cultivation of applied talents as their main goal, emphasizing students' practical ability and application ability. Therefore young teachers need to continuously improve their abilities in teacher ethics, professional knowledge, pre-teaching preparation, communication and cooperation. The development of teachers' competence is a process of long-term continuous change and influenced by multiple factors, which coincides with the concept of Dynamic Systems Theory (DST). The progress of young teachers' competence in private applied colleges and universities is not achieved overnight, but is gradual and dynamic, so there should be a scientific order for learning various aspects of competence in order to maximize teaching competence and level in a short period of time. In this regard, this paper combines reinforcement learning theory and knowledge mapping to design a recommendation algorithm for teacher ability improvement.

2.1. Dynamic Systems Theory

Dynamic systems (DST) theory provides us with a scientific perspective for portraying the dynamic, complex, and nonlinear characteristics presented in the growth process of young teachers in private applied colleges and universities. The system is characterized by the intertwined roles of operational entities and participants. The overall structure is built by a series of embedded subsystems, and there is a deep network of interaction among these subsystems. According to the theory of dynamic systems, the term "dynamic" refers to the evolution of a system over time, which may be driven by internal or external energy, while "complex" refers to the interaction between the system and its subsystems that leads to changes, which are often sudden and discontinuous. "Nonlinear", on the other hand, means that the evolution of a system cannot be understood purely through linear causality, and it can be variable, disordered, and unpredictable. DST theory holds that whether it is a system, a subsystem, or a subsystem, they all follow similar working principles at different levels.

2.2. Deep reinforcement learning

In deep reinforcement learning, two basic concepts, namely the intelligent body and the environment, together constitute the basic framework of the deep reinforcement learning problem [23]. Intelligent body, i.e., the subject that performs the learning task, its main purpose is to learn the strategy through continuous interaction with the environment, and to be able to select the optimal action in different states in order to obtain the maximum cumulative reward. The environment is the entire external system except the intelligent body, and the environment will change accordingly to the actions chosen by the intelligent body. The interaction process between the intelligent body and the environment is shown in Figure 1. The figure shows the process of constant interaction between the intelligent body and the environment. The intelligent body gives action a according to the initial state s of the user (i.e., the environment), and after executing action a , the user's state changes to s' , and the reward generated after executing action a is given to r' , and the intelligent body decides the next action according to the strategy a' until the end of the interaction.

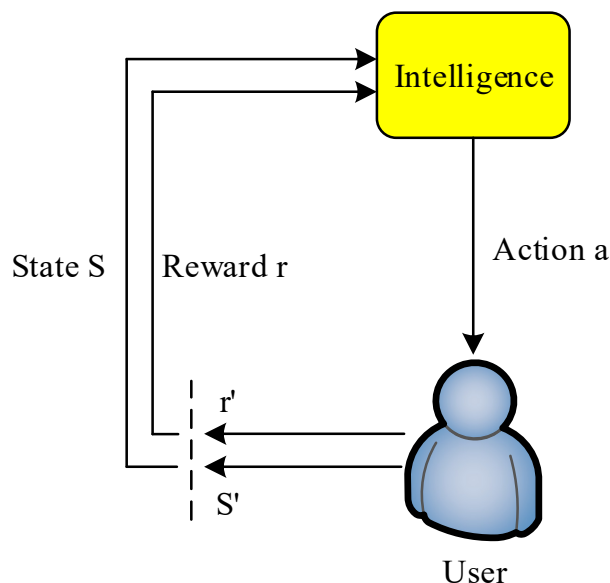


Fig. 1. Intelligent body and environmental interaction process

The mathematical basis of the interaction process between an intelligent body and its environment is the Markov Decision Process (MDP), which consists of a set of five tuples $\langle S, A, R, P, \gamma \rangle$, which need

to be adjusted for different specific problems, and the significance of each element of the set of five tuples is described in detail in the following:

State space S : The state space is the set of all possible states of the environment, and state s is an instantiation of the state space, which can be either discrete or continuous. At some point in time, the environment will be in one of the states in the state space, and the intelligence will make decisions based on the current state.

Action Space A : The action space contains all the possible actions that an intelligent body can perform. These actions can be discrete or continuous. In a certain state, the environment will provide a set of all possible actions that can be executed, and the intelligent body will choose to execute one of the actions in the action space to change the environment according to the current strategy, and the environment will change the state at this moment after receiving the rewards generated by the actions.

Reward Function R : The reward function defines the immediate reward that an intelligent body receives for performing an action in a given state. Rewards can be real, immediate, delayed, positive or negative, depending on the problem. The reward function guides the learning process of the intelligent body so that it can learn to take optimal actions in different states.

State Transfer Probability P : The state transfer probability describes the probability distribution of the environment transferring to the next state after the intelligent body performs an action. Specifically, for each state and action pair in the state space, the state transfer probability function gives the probability that the next state is each of the possible states after the execution of the action.

Discount Factor γ : The discount factor is used to measure the importance of future rewards, i.e., the extent to which their influence on an intelligent's action choice can be suppressed by multiplying them with future rewards. In general $\gamma \in [0, 1]$, when the discount factor is close to 1, the intelligence pays more attention to future rewards. When the discount factor is close to 0, the intelligent body pays more attention to immediate rewards and gives a larger discount to future rewards the further in the future it is from the current state.

2.3. Knowledge mapping

A knowledge graph is a knowledge base consisting of a number of entities with different attributes, and different relationships between different entities in a mesh [24]. From a graphical point of view, a knowledge graph compares people and objects in the real world to individual entities in a graph, and their relationships in the real world to edges between entities, so it seems that a knowledge graph is also a symbolic representation of real time, a semantic network that maps the objective world in a clearer and more concise way, and expresses the connections of the real world in terms of simple nodes and edges. The knowledge graph can be understood as a heterogeneous information network containing entities (nodes) and relations (directed edges), denoted by $G = \{V, E\}$. $V = \{P\}$ represents the set of entities in the knowledge graph and E represents the set of edges.

A tuple is a common representation of a knowledge graph, commonly described as a ternary structure in the form of e.g. <Head Entity, Relationship, Tail Entity> or <Entity, Attribute, Attribute Value>, where an entity represents a unique distinguishable something that exists independently. The ternary is formalized as: $G = \{(h, r, t) | h, t \in \mathcal{E}, r \in \mathcal{R}\}$, where $\mathcal{E}$ is the set of entities, $\mathcal{R}$ is the set of relations, h is the head entity, t is the tail entity, and r is the relation between the two of them. Knowledge graphs applied to the recommendation domain can significantly improve the performance of recommendations and alleviate the data sparsity and cold-start problems in recommendations.

3. Recommended algorithms for teacher empowerment

This chapter proposes the Knowledge Graph Teacher Competency Enhancement Recommendation Algorithm (KGDR) based on Deep Reinforcement Learning, which is applied to the dynamic development mechanism of teacher competency.

KGDR contains three parts: young teacher competency knowledge graph module, attention deep reinforcement learning module and path reasoning module. Among them, the young teacher competency knowledge mapping module contains the vector representation of entities and relationships, the path information of knowledge mapping and the reward function of reinforcement learning. The Attention Deep Reinforcement Learning module passes the input state information and action space information through the attention layer and the four fully connected layers, and finally obtains the value network $\hat{v}(s_t)$ and the strategy network $\pi(\cdot|s_t, \tilde{A}_{u,t})$, of which the value network $\hat{v}(s_t)$ will be used as the value baseline for reinforcement learning, and the strategy network $\pi(\cdot|s_t, \tilde{A}_{u,t})$ will be used in the path inference module for the inference of the path. The Random_Beam_Search search algorithm proposed in this paper is used in the path inference module to search and sort to get the final recommendation results and inference paths.

3.1. Young Teacher Competency Knowledge Mapping Module

The Markov decision process introduced in this module allows the Young Teacher Competency Knowledge Graph to be incorporated as a dataset for the recommendation algorithm [25]. The ternary of state s_t of a node at a moment t of the Young Teacher Competency Knowledge Graph is defined as (u, e_t, h_t) , where u represents the initial Young Teacher entity, e_t represents the entity reached by the intelligent body's reasoning at moment t , h_t represents the reasoning path before moment t which can be represented as $\{u_0, r_1, e_1, \dots, r_t, e_t\}$, s_0 of the initial state is defined as $(u, u, \emptyset)$, and s_T of the final state is defined as (u, e_T, h_T) . The behavioral space A_t of this intelligent body represents all connected edges e_t at the moment of t time, but does not include entities and relationships of e_t history.

In this paper, the pruning strategy is introduced to solve the inefficiency problem of exploring intelligences due to large data, and the specific formula of the pruning strategy is shown as follows:

$$\tilde{A}_t(u) = \{(r, e) | \text{rank}(f((r, e) | u))\} < \alpha, (r, e) \in A_t \quad (1)$$

Where α denotes the maximum action space in which the intelligent body can carry out exploration, and the value needs to be set manually, and function $f((r, e) | u)$ is a kind of multi-hop scoring function.

Introducing the concept of multi-hop path into the knowledge graph module of young teacher competency, assuming that there exists a certain node e_i in the middle of the path of the knowledge graph, i.e.: $\left\{e_0 \xleftrightarrow{r_1} e_1 \xleftrightarrow{r_2} \dots e_i \xleftrightarrow{r_{i+1}} \dots e_k\right\}$, the multi-hop scoring function can be defined as:

$$f(e_0, e_k | \tilde{T}_{i,k}) = \left\langle e_0 + \sum_{s=1}^i r_s, \sum_{s=i+1}^k r_s + e_k \right\rangle + b_{ek} \quad (2)$$

Where $\langle \dots \rangle$ represents the dot product operation, i.e., the vectors composed of e_0 to e_i are multiplied by the vectors composed of e_i to e_k , and e, r represents the multidimensional vectors of entities and relations, respectively, and the function $f((r, e) | u)$ in Eq. (1) can be computed using $f((u, e) \tilde{r}_{i,k})$.

Reward function: In this paper, the reward function R_t of reinforcement learning is expressed as a relative similarity between young teachers u and teacher competence Z_T when and only when teacher competence Z_T belongs to the set of recommended goals after pruning in action space A_t , and its formula is:

$$R_t = \begin{cases} \max \left(0, \frac{f(u, Z_T)}{\max_{i \in E} f(u, i)} \right) & \text{if } Z_T \in E \\ 0 & \text{otherwise} \end{cases} \quad (3)$$

where E denotes all teacher competencies in the knowledge graph, $f(u, i)$ denotes the scoring function between u and any teacher competency i two nodes, and $f(u, Z_T)$ denotes the scoring function between u , Z_T two nodes. Both functions are derived from the one-hop scoring function of Eq. (2) The specific definition of the one-hop scoring function can be expressed as follows: $f(u, e | \tilde{r}_{1,1})$.

3.2. Attention Deep Reinforcement Learning Module

The specific definition of each layer of the network in this module is shown below.

The state s_t is taken as input, and the input x_1 of the next module is obtained after full connection, activation function and dropout operations in turn, see equation (4):

$$x_1 = \text{dropou}(\sigma(s_t W_1)) \quad (4)$$

where the activation function σ represents the ReLU function.

The result obtained by passing x_1 through the attention layer is passed through the activation function $\tanh$ to obtain the attention weights attention_weights , see equation (5):

$$\text{attention_weights} = \tanh(x_1 \cdot \text{Attention}) \quad (5)$$

Putting attention_weights through the fully connected layer and softmax activation operation yields the probability distribution result x_2 , see equation (6). Multiply x_2 with x_1 matrix to get the result x with attention mechanism, see equation (7):

$$x_2 = \text{soft max}(\text{attention_weights} \cdot W_2) \quad (6)$$

$$x = x_1 \cdot x_2 \quad (7)$$

The result obtained from x in Eq. (7) by fully connected layer operation is hadamard product with the action space $\tilde{A}_t(u)$ and the result is subjected to softmax activation operation to obtain the policy network, see Eq. (8):

$$\pi(\cdot | s_t, \tilde{A}_{u,t}) = \text{soft max}(\tilde{A}_u \otimes (x W_p)) \quad (8)$$

where $\otimes$ denotes the Hadamard product, an operation that shields illegal actions outside the action space so that all actions are in the action space $\tilde{A}_t(u)$.

The value network $\hat{v}(s_t)$ is obtained by passing x in Eq. (7) through the fully connected layer operation, see Eq. (9):

$$\hat{v}(s_t) = x W_v \quad (9)$$

$$\theta = (\text{Attention}, W_1, W_2, W_p, W_v) \quad (10)$$

The ultimate goal of the algorithm proposed in this paper is to learn a stochastic strategy π such that for any initial young teacher its maximized cumulative reward is obtained, expressed as follows:

$$R(\theta) = E_{\pi} \left[\sum_{t=0}^{T-1} \gamma^t R_{t+1} | s_0 = (c, c, \emptyset) \right] \quad (11)$$

where γ denotes the discount factor and R_{t+1} denotes the reward value received at moment t .

In order to allow the intelligent body in reinforcement learning to explore more paths, this paper introduces the concept of regularization to maximize θ , which is defined as follows:

$$\nabla_{\theta} R(\theta) = E_{\pi} \left[\nabla_{\theta} \log \pi_{\theta}(\cdot | s, \tilde{A}_u) (R - \hat{v}(s)) \right] \quad (12)$$

The purpose of introducing regularization is to avoid overfitting caused by too large a parameter, where R represents a cumulative reward from state S to state S_T , and $\hat{v}(s)$ represents the value network obtained, the cumulative reward minus the current value can be obtained from the value of its historical path, which is conducive to improving the efficiency of the model to find the optimal goal.

3.3. Path Reasoning Module

The Random_Beam_Search algorithm introduced in the path inference module of this paper is an improvement of the Beam_Search algorithm, which can expand the search space to avoid falling into the local optimum during the search process to improve the possibility of the algorithm to find the optimal solution. The inputs of the Random_Beam_Search algorithm are the young teacher u , the final moments T and $top\ k = \{K_1 \cdots K_T\}$, and the final outputs are the recommended path P_T , probability value Q_T , and reward value R_T . Find all the paths corresponding to each moment and the action space of the target young teacher after the pruning operation, respectively, based on the action space to find the probability set of each search path after different search actions at the current moment $p(\alpha)$. For each search space K in $top\ k = \{K_1 \cdots K_T\}$ is expanded to make it 1.3 times the original, the search algorithm randomly selects K values from the expanded search space, and selects the top K search paths from the set of path probabilities $p(\alpha)$ as the recommended alternative paths. Each of the selected search paths is processed in turn to save the obtained probability values, reward values and recommended paths into Q_T , R_T and P_T , and finally the Random_Beam_Search algorithm outputs the final Q_T , R_T , and P_T as a way to realize the provision of inference paths for the recommended results and to improve their diversity.

4. Empirical analysis

4.1. Model comparison experiments

4.1.1. Evaluation indicators. KGDR is evaluated using four main indicators.

In this experiment, accuracy is defined as the ratio of the number of correct predictions to the overall predictions, as shown in Equation (13):

$$\text{Precision} = \frac{TP}{TP + FP} \quad (13)$$

Recall is defined as the proportion of all true positive cases that are correctly predicted as positive by the model, as shown in Equation (14):

$$\text{Recall} = \frac{TP}{TP + FN} \quad (14)$$

The F1 score is a reconciled average of correctness and recall and provides a comprehensive performance metric. The calculation formula (15) is shown:

$$F_1 = \frac{2TP}{2TP + FP + FN} \quad (15)$$

The F1 score takes into account the trade-off between correctness and recall. It ranges from 0 to 1, with higher scores indicating that the model performs better in balancing precision and recall.

Finally, the diversity metrics are defined using Equation (16):

$$Diversity = \frac{1}{|U|} \sum_{u \in U} \frac{\sigma_{domains}(R(u))}{\mu_{domains}(R(u))} \quad (16)$$

Where U is the set of capability enhancement, $R(u)$ is the recommendation list of capability enhancement u , $\sigma_{domains}$ is the standard deviation of different capability enhancement priorities in the recommendation list, and $\mu_{domains}$ is the average of different capability enhancement priorities. The diversity metric evaluates the adaptability of the recommender system in that the larger the metric, the better the diversity of the recommender system in terms of providing ability boosting paths.

4.1.2. Experimental results. The SVD, LFM, NCF, LinUCB, TS, DBGD, and MARL algorithmic models were selected for the study to compare their recommendation effectiveness with the KGDR model on the Coursera dataset, respectively. The top- k recommendation was performed when comparing the KGDR model with other algorithmic models.

The results of model recall, accuracy, F1 and diversity comparison are shown in Figures 2 to 5, respectively. A single recommendation usually leads to a significant decrease in the metrics. When $k = 20$, these three metrics are usually higher than when $k = 10$ or $k = 30$, which suggests that the recommendation list at $k = 20$ is more likely to enhance the competence of young teachers. The KGDR model, compared to other models, at $k = 20$ The values of all indicators are the highest, which indicates that the KGDR model has better recommendation effect.

The results of model diversity comparison show that this paper's model at $k = 20$, $Diversity = 0.7876$, are higher than the other algorithmic models, indicating that the KGDR model can provide more diverse paths for teacher competence improvement.

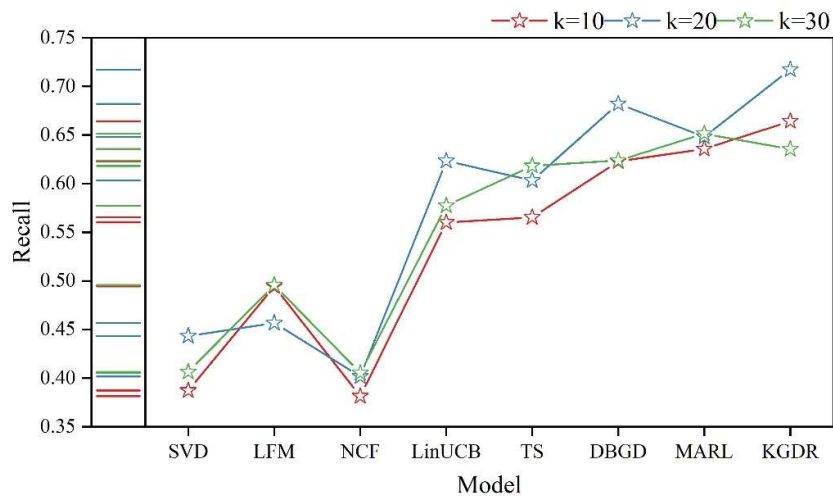


Fig. 2. Recall comparison

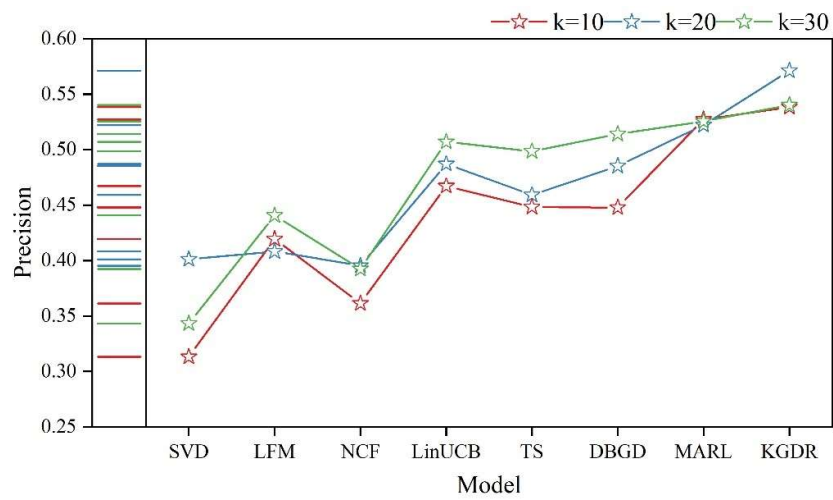


Fig. 3. Accuracy comparison

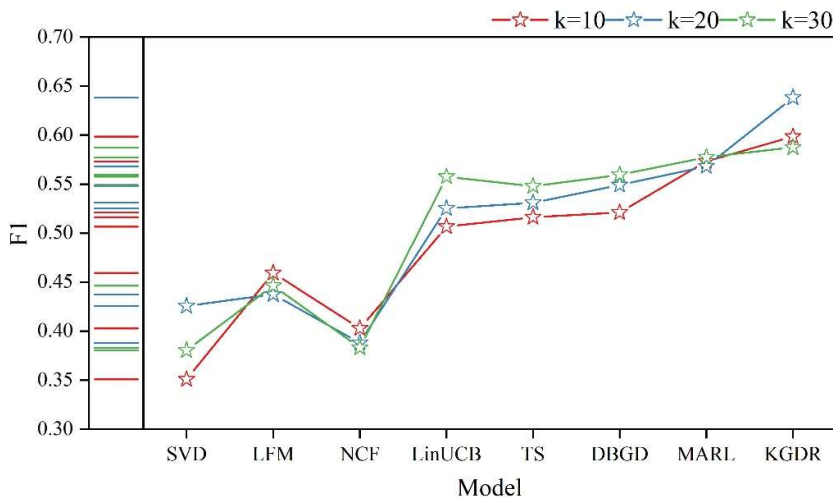


Fig. 4. F1 comparison

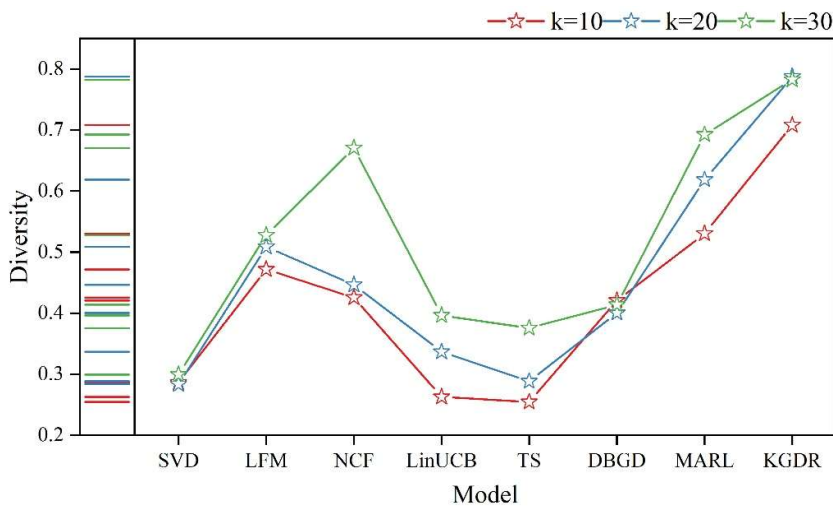


Fig. 5. Diversity comparison

4.2. *Analysis of the effectiveness of the application of dynamic development mechanisms*

In this section, 30 young teachers employed in the first semester of 2023 in a private applied university in Y city were used as the research subjects. The dynamic development mechanism of young teachers' competence based on KGDR model designed in this paper was used for training. By the end of the next semester in 2024, the performance evaluation based on teacher competency improvement was conducted for the 30 young teachers using the fuzzy comprehensive evaluation method.

4.2.1. **Weights of Performance Evaluation Indicator System for Young Teachers.** The index system and weight setting of young teachers' performance evaluation based on teachers' ability improvement in the research reference [26] are shown in Table 1. This chapter utilizes this evaluation index system to explore the effect of the application of the dynamic development mechanism based on KGDR. The whole evaluation index system has three levels of indicators. The first-level indicators are each of the indicators at the first level within the structure of the system, which are dimensionally divided into four aspects: teacher ethics, teaching practice, comprehensive parenting, and personal development. Level 2 indicators are the indicators at the second level within the structure of the evaluation indicator system, reflecting the aspects subordinate to the level 1 indicators, including 11 indicators of teacher ethics, educational sentiment, professional knowledge, pre-teaching preparation, curriculum teaching, teaching results, total work, class guidance, curriculum nurturing, professional growth, and communication and cooperation. The third-level indicators in turn use the second-level indicators as the object of evaluation.

Table 1. The performance evaluation index system and weight setting of teachers

Primary index	W_i	Secondary index	W_i	Tertiary index	W_i	Composite weight
Master teacher (A)	0.352	Standard (A1)	0.257	Ideal belief (A11)	0.75	0.193
				Nurturing Talents through Moral Education (A12)	0.25	0.064
		Education (A2)	0.095	Teaching by aptitude (A21)	0.78	0.074
				Self-cultivation (A22)	0.22	0.021
Teaching practice (B)	0.332	Professional knowledge (B1)	0.075	Education foundation (B11)	0.65	0.049
				Subject quality (B12)	0.264	0.02
				Information literacy (B13)	0.086	0.006
		Pre-teaching preparation (B2)	0.036	Familiar with course (B21)	0.234	0.008
				Teaching skill (B22)	0.411	0.015
				Design teaching (B23)	0.355	0.013
		Course teaching (B3)	0.051	Teaching group (B31)	0.77	0.039
				Teaching evaluation (B32)	0.23	0.012
		Teaching results (B4)	0.146	Student achievements and honors (B41)	0.314	0.046
				Students' daily (B42)	0.129	0.019
				Students' physical and mental health (B43)	0.557	0.081
		Total amount of work (B5)	0.024	Teaching time (B51)	0.64	0.015
				Teaching activity (B52)	0.36	0.009
Holistic (C)	0.105	Class guidance (C1)	0.026	Class management (C11)	0.31	0.008
				Psychological counseling (C12)	0.69	0.018
		Course education (C2)	0.079	Childcare practice (C21)	0.46	0.036
				Finishing course (C22)	0.54	0.043
Personal development (D)	0.211	Professional growth (D1)	0.07	Development planning (D11)	0.46	0.032
				Reflective improvement (D12)	0.26	0.018
				Honorable mention (D13)	0.28	0.02
		Communication cooperation (D2)	0.141	Communication skills (D21)	0.61	0.086
				Team collaboration (D22)	0.39	0.055

4.2.2. Analysis of the effect of improving the capacity of young teachers. Before the start of the experiment, 10 experts were invited to evaluate the performance of 30 young teachers in terms of

teacher competence according to the evaluation index system. The evaluation grades are divided into 4 grades: excellent, good, moderate and poor. The statistical evaluation set is shown in Table 2. According to the weights of the indicators and the evaluation set, the fuzzy comprehensive rating method is used to calculate the affiliation value of $[0, 0.195, 0.201, 0.604]$. According to the principle of maximum affiliation, 0.604 is the largest, corresponding to the evaluation grade of “poor”.

Table 2. Young teachers ability to improve evaluation set

Tertiary index	Excellence	Good	Medium	Bad	Tertiary index	Excellence	Good	Medium	Bad
A11	0	0.2	0.2	0.6	B42	0	0.3	0.3	0.4
A12	0	0.2	0.2	0.6	B43	0	0.3	0.2	0.5
A21	0	0.2	0.2	0.6	B51	0	0.2	0.2	0.6
A22	0	0.2	0.2	0.6	B52	0	0.2	0.2	0.6
B11	0	0.2	0.1	0.7	C11	0	0.1	0.3	0.6
B12	0	0.1	0.2	0.7	C12	0	0.1	0.2	0.7
B13	0	0.2	0.2	0.6	C21	0	0.2	0.2	0.6
B21	0	0.1	0.1	0.8	C22	0	0.2	0.2	0.6
B22	0	0.2	0.2	0.6	D11	0	0.2	0.2	0.6
B23	0	0.2	0.2	0.6	D12	0	0	0.5	0.5
B31	0	0.2	0.2	0.6	D13	0	0.1	0.2	0.7
B32	0	0.2	0.1	0.7	D21	0	0.2	0.2	0.6
B41	0	0.1	0.2	0.7	D22	0	0.2	0.2	0.6

After one semester of training using the mechanism of dynamic development of young teachers' competence based on the KGDR model, 10 experts were invited to evaluate the performance again. The statistical set of evaluations is shown in Table 3. The number of “poor” ratings is zero for all 26 evaluation indicators.

Table 3. Young teachers ability to improve evaluation set

Tertiary index	Excellence	Good	Medium	Bad	Tertiary index	Excellence	Good	Medium	Bad
A11	0.6	0.4	0	0	B42	0.4	0.5	0.1	0
A12	0.4	0.3	0.3	0	B43	0.5	0.3	0.2	0
A21	0.4	0.3	0.3	0	B51	0.6	0.4	0	0
A22	0.5	0.3	0.2	0	B52	0.4	0.3	0.3	0
B11	0.6	0.3	0.1	0	C11	0.6	0.2	0.2	0
B12	0.7	0.2	0.1	0	C12	0.7	0.2	0.1	0
B13	0.6	0.3	0.1	0	C21	0.6	0.3	0.1	0
B21	0.6	0.2	0.2	0	C22	0.5	0.4	0.1	0
B22	0.6	0.3	0.1	0	D11	0.4	0.4	0.2	0
B23	0.5	0.4	0.1	0	D12	0.5	0.3	0.2	0
B31	0.6	0.3	0.1	0	D13	0.7	0.2	0.1	0
B32	0.5	0.4	0.1	0	D21	0.6	0.1	0.3	0
B41	0.4	0.4	0.2	0	D22	0.6	0.4	0	0

The affiliation value calculated using the fuzzy composite rating method is $[0.54, 0.32, 0.14, 0]$, and according to the principle of maximum affiliation, 0.54 is the maximum, which corresponds to an evaluation grade of “excellent”. This evaluation result is two grades higher than that before the implementation of the dynamic development mechanism. This shows that the dynamic development mechanism of young teachers' competence based on KGDR model designed in this paper can effectively improve the competence of young teachers.

5. Conclusion

The study designed a deep reinforcement learning-based knowledge graph teacher competence improvement recommendation algorithm to recommend relevant paths for the dynamic development of young teachers' competence. Compared with other single recommendation models, this paper's model has the highest values of recall, accuracy, F1 and diversity evaluation indexes, indicating that this paper's model works better for young teachers' competence improvement. Thirty young teachers from a private applied university in city Y were used as the research subjects. After one semester of using the model, the fuzzy comprehensive evaluation grade of young teachers is "excellent". This study can provide theoretical value and practical reference for the optimization of young teachers' competence development mechanism in other colleges and universities. Future research can further explore other application areas of intensive learning in the development of teachers' competence, such as the enhancement of teachers' teamwork ability, in order to realize the comprehensive and coordinated development of teachers' competence.

References

- [1] Liu, L. (2023). Thoughts on the Application-oriented Private Universities' Implementation of Practical Teaching through Industry-Education Integration. *Journal of Human Resource Development*, 5(4), 70-77
- [2] Zhang, H. (2023). Construction of Application-oriented Practical Teaching System in Private Undergraduate Colleges Based on OBE Concept. *International Journal of Education and Humanities*, 8(3), 213-219.
- [3] Chen, M. (2020, December). Advantages of Operational Mechanism of Application-Oriented Private Colleges—Taking Nantong Institute of Technology as an Example. In *2020 6th International Conference on Social Science and Higher Education (ICSSHE 2020)* (pp. 169-173). Atlantis Press.
- [4] Zhang, L. (2022). Strategy and Path Innovation for High-Quality Development of First-Class Private Application-Oriented Universities. *Adult and Higher Education*, 4(7), 86-92.
- [5] Lu, Y. (2021). Research and analysis on the improvement of teaching ability of young teachers in colleges and universities. *J. High. Educ. Res*, 2(6), 329-336.
- [6] Guo, Q., He, Y., Zhang, J., Liu, Z., Peng, J., & Li, L. (2024). Assessing and Enhancing the Professional Development Capacity of Young Teachers in Universities and Colleges: Differences Between Young Teachers in China and the United States. *Journal of Humanities, Arts and Social Science*, 8(7).
- [7] Wang, N. (2019). Research on the Development Strategy of Young Teachers in Vocational Colleges under the Background of Professional Development. In *2019 9th International Conference on Education and Management (ICEM 2019)* (pp. 399-407).
- [8] Peng, F., Zhu, Q., & Wang, S. (2023). Study on the Career Development of Young Teachers in Vocational Colleges. *Journal of Contemporary Educational Research*, 7(8), 11-17.
- [9] Liu, J., & Zhang, Y. (2017). Implementation of Information-based Teaching System for Young College Teachers Based on iOS Platform. *International Journal of Emerging Technologies in Learning*, 12(8).
- [10] Arulkumaran, K., Deisenroth, M. P., Brundage, M., & Bharath, A. A. (2017). Deep reinforcement learning: A brief survey. *IEEE Signal Processing Magazine*, 34(6), 26-38.
- [11] Ernst, D., & Louette, A. (2024). Introduction to reinforcement learning. Feuerriegel, S., Hartmann, J., Janiesch, C., and Zschech, P, 111-126.
- [12] Toom, A. (2017). Teachers' professional and pedagogical competencies: A complex divide between teacher work, teacher knowledge and teacher education. *The SAGE handbook of research on teacher*

education, 2, 803-819.

- [13] Kleickmann, T., Richter, D., Kunter, M., Elsner, J., Besser, M., Krauss, S., & Baumert, J. (2013). Teachers' content knowledge and pedagogical content knowledge: The role of structural differences in teacher education. *Journal of teacher education*, 64(1), 90-106.
- [14] Widana, I. W. (2020, July). The effect of digital literacy on the ability of teachers to develop HOTS-based assessment. In *Journal of Physics: Conference Series* (Vol. 1503, No. 1, p. 012045). IOP Publishing.
- [15] Basma, B., & Savage, R. (2018). Teacher professional development and student literacy growth: A systematic review and meta-analysis. *Educational psychology review*, 30, 457-481.
- [16] Nalipay, M. J. N., Mordeno, I. G., Semilla, J. R. B., & Frondoza, C. E. (2019). Implicit beliefs about teaching ability, teacher emotions, and teaching satisfaction. *The Asia-Pacific Education Researcher*, 28, 313-325.
- [17] Rahmawati, I., Ghifariand, R. F., & Lestari, H. (2020). Enhancing the Effectiveness of Teacher Work and Teams. *KnE Social Sciences*, 484-492.
- [18] Du, J. (2021). Research and practice of training effect evaluation system based on Kirkpatrick model—taking teacher training system of Zhejiang open university as an example. *Chinese Studies*, 10(2), 123-146.
- [19] Loyalka, P., Popova, A., Li, G., & Shi, Z. (2019). Does teacher training actually work? Evidence from a large-scale randomized evaluation of a national teacher training program. *American Economic Journal: Applied Economics*, 11(3), 128-154.
- [20] Avidov-Ungar, O. (2016). A model of professional development: Teachers' perceptions of their professional development. *Teachers and teaching*, 22(6), 653-669.
- [21] Micheeva, T., Murugova, E., Morozova, Y., & Demchenko, V. (2020). Training as a major tool of teacher professionalism enhancement. In *INTED2020 Proceedings* (pp. 1211-1215). IATED.
- [22] Sayed, Y., & Bulgrin, E. (2020). Teacher professional development and curriculum: Enhancing teacher professionalism in Africa. *Education International*.
- [23] Jaime Govea, Alexandra Maldonado Navarro, Santiago Sánchez Viteri & William Villegas Ch. (2024). Implementation of deep reinforcement learning models for emotion detection and personalization of learning in hybrid educational environments. *Frontiers in Artificial Intelligence* 1458230-1458230.
- [24] Zhihao Guo, Peng Song, Chenjiao Feng, Kaixuan Yao, Chuangyin Dang & Jiye Liang. (2024). Causal intervention for knowledge graph denoising in recommender systems. *International Journal of Machine Learning and Cybernetics*(prepublish), 1-17.
- [25] Cyrille Vessaire, Pierre Carpentier, Jean Philippe Chancelier, Michel De Lara & Alejandro Rodríguez Martínez. (2024). Complexity Bounds for Deterministic Partially Observed Markov Decision Processes. *Annals of Operations Research*(prepublish), 1-38.
- [26] Wang Lili, Wen Hengfu & Liu Yang. (2016). AHP Based Quantitative Evaluation Index System of Teacher's Research Performance in the University. *International Journal of Multimedia and Ubiquitous Engineering*(7), 391-402.