

A Knowledge Graph Reasoning-Based Model for Computerized Adaptive Testing

Xinyi Qiu

East China Normal University
Shanghai, China
xyqiu@stu.ecnu.edu.cn

Zhiyun Chen *

East China Normal University
Shanghai, China
chenzhy@cc.ecnu.edu.cn

Abstract

The significant of Computerized Adaptive Testing (CAT) is self-evident in contemporary Intelligent Tutoring Systems (ITSs) which aims to recommend suitable questions for students based on their knowledge state. In recent years, Graph Neural Networks (GNNs) and Reinforcement Learning (RL) methods have been increasingly applied to CAT. While these approaches have achieved empirical success, they still face limitations, such as inadequate handling of concept relevance when multiple concepts are involved and incomplete evaluation metrics. To address these issues, we propose a Knowledge Graph Reasoning-Based Model for CAT (KG-CAT), which leverages the reasoning power of knowledge graphs (KGs) to capture the semantic and relational information between concepts and questions while focusing on reducing the noise caused by concepts with low relevance by utilizing mutual information. Additionally, a multi-objective reinforcement learning framework is employed to incorporate multiple evaluation objectives, further refining question selection and improving the overall effectiveness of CAT. Empirical evaluations conducted on three authentic educational datasets demonstrate that the proposed model outperforms existing methods in both accuracy and interpretability.

1 Introduction

In recent years, with the rapid development of computer technology, ITSs have been widely applied in the field of education. Compared with the traditional paper test, CAT has become a common means of modern examination with its remarkable advantages of automation and personalization (Weiss and Kingsbury, 1984; Chen et al., 2015). As illustrated in Figure 1(a), CAT primarily consists of two main components: Cognitive Diagnosis Model (CDM) and Selection Algorithm (SA). CDM utilizes statistical or machine learning methods to

estimate the current knowledge level of students based on their historical interaction data. The most commonly used CDMs include Item Response Theory (IRT) (Embretson and Reise, 2013) and Neural Cognitive Diagnosis (NCD) (Wang et al., 2020a), which are based on psychology and deep learning respectively.

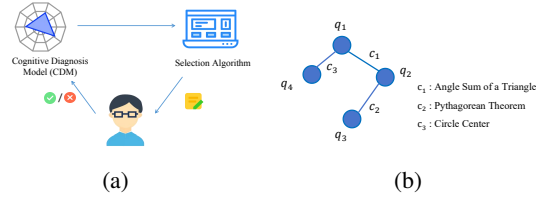


Figure 1: (a) The components of CAT. (b) An example of the relationship between questions and concepts.

In recent years, GNNs (Scarselli et al., 2009) have demonstrated exceptional performance across various fields, leading to their application in CAT as well. (Nakagawa et al., 2019) represented concepts and their relationships using nodes and edges, but failed to account for distinctions between individual questions associated with the same concept. (Yang et al., 2021) and (Wang et al., 2023) incorporated questions into the graph structure by using nodes to represent both the concepts and the questions, with an edge connecting each question to its corresponding concept and successfully solved the above problem. However, these approaches overlook the natural distinction between concepts and questions. Furthermore, existing studies have failed to consider the potential impact of varying concept relevance when a question involves multiple concepts during the training process. The differing significance of each concept in relation to the question may influence the overall learning outcomes, which has not been adequately addressed in prior work. Figure 1(b) is an example, q_1 contains concept *Angle Sum of a Triangle* and *Circle Center*, q_2 contains concepts *Angle Sum of*

*Corresponding author

a *Triangle* and *Pythagorean Theorem*, q_3 contains concept *Pythagorean Theorem* and q_4 contains concept *Circle Center*. When predicting a student’s probability of correctly answering question q_2 , it is essential to leverage the historical response data from questions related to q_2 . In this context, it is more appropriate to prioritize information from q_3 over q_1 , as q_1 involves additional concepts, such as *Circle*, which may introduce noise into the prediction process. Meanwhile, previous studies (Bi et al., 2020; Zhuang et al., 2022; Wang et al., 2023) have only focused on limited objectives, neglecting the important influence of question difficulty and student similarity.

To address these issues, we propose a model based on knowledge graph reasoning, which represents questions as nodes and concepts as edges to better capture the semantic and relationship information between them. To specifically reduce the noise introduced by irrelevant concepts, we introduce a disentanglement module that utilizes mutual information. This module enhances the independence of various concepts and helps mitigate the impact of unrelated concepts on the prediction process. Furthermore, acknowledging that existing selection algorithms inadequately consider question difficulty and student similarity, we integrate a multi-objective reinforcement learning framework. This addition ensures more accurate recommendations by incorporating these factors into the question selection process.

In summary, the main contributions are as follows:

- We propose a Knowledge Graph Reasoning-Based Model for Computerized Adaptive Testing (KGCAT) to better capturing the relationships and semantic information between questions and concepts while focus on introducing mutual information to disentangle multiple concepts and reduce the information interference of unnecessary neighbors when aggregating nodes.
- We employ a multi-objective reinforcement learning framework and incorporate new considerations of question difficulty and student similarity to optimize the overall performance of the CAT model.
- We conduct extensive experiments on public datasets to demonstrate that KGCAT outperforms existing methods in terms of accuracy

and interpretability.

2 Related Work

2.1 Computerized Adaptive Testing

In recent years, the improvement of CAT can be mainly divided into two categories : CDM-based models (Tong et al., 2022; Yang et al., 2021; Nakagawa et al., 2019) and SA-based models (Zhuang et al., 2022; Wang et al., 2023; Bi et al., 2020; Ghosh and Lan, 2021). (Yang et al., 2021) and (Nakagawa et al., 2019) proposed the use of GNNs to capture the relationships between concepts and questions, guiding the model’s learning process. On the other hand, (Tong et al., 2022) suggested that in addition to the explicit relationships between questions and concepts, the implicit connections, derived from students’ responses, can also be leveraged to better structure these relationships. In contrast to these approaches, our method uses a knowledge graph to represent both concepts and questions, simultaneously capturing their semantic and relational information.

For the SA-based approaches, (Bi et al., 2020) proposed a model-independent framework that considers quality and diversity to support multiple CAT scenarios. (Zhuang et al., 2022) and (Wang et al., 2023) transformed CAT into a reinforcement learning problem and introduces similar metrics to guide question selection. Building upon the metrics established in prior research, we introduce novel evaluation criteria to further guide the question selection process.

2.2 Knowledge Graph

Knowledge Graphs (KGs) have become widely adopted in recommendation systems due to their ability to effectively capture and represent relationships between entities. Existing KG-based recommendation models can be broadly categorized into three categories: embedding-based (Gao et al., 2022; Cao et al., 2019), path-based (Catherine and Cohen, 2016; Wang et al., 2019b), and GNN-based (Wang et al., 2019a; Ai et al., 2018).

The embedding-based methods use KG embedding techniques to construct knowledge graphs. The path-based methods extract the paths from the users to the items through the KG entities, and input into RNN and memory network to guide the recommendation. GNN-based methods use GNN to learn long-range information through propagation. For example, (Wang et al., 2019a) explicitly

modelled and captured high-order relationships between users and items in the GNN framework and (Ai et al., 2018) constructed a knowledge graph at a fine-grained intention level and preserved semantics using relational dependencies. However, the application of knowledge graphs in CAT scenarios has not been fully explored. Therefore, our work will focus on investigating the potential and applications of knowledge graphs in the CAT context.

3 Method

3.1 The structure of KG-CAT

Figure 2 is an overall structure framework of KG-CAT, which mainly includes two modules: KGSE (Knowledge Graph-Based State Encoder) and DSMORL (Difficulty and Similarity-Based Multi-Objective Reinforcement Learning). KGSE encodes the student state s_{t-1} input at time t into $\tilde{s}_t$ and outputs it to DSMORL. DSMORL uses RL to select the question q_t , which will be combined with the concept c_t and response y_t at time $t + 1$ as the input s_t at the next time.

3.2 KGSE

A question may involve multiple concepts, and sometimes there is a certain overlap between these concepts. In order to reduce the coupling between these concepts, a question should be separated into independent variables. Therefore, the feature of a question is projected onto K latent spaces:

$$q_i = \{h_{i,1}, h_{i,2}, \dots, h_{i,K}\}, h_{i,K} \in \mathbb{R}^{\frac{d_{embed}}{K}} \quad (1)$$

where $h_{i,K}^0 = \sigma(W_K \bullet x_i)$ is the initial embedding, W_K is the K -th projection matrix, x_i is the feature vector and σ is the activation function.

Since the information carried by an entity is often incomplete, it is necessary to aggregate neighbor nodes to learn richer features. However, not all the information of neighbor nodes need to be learned, the model should learn the neighbor nodes related to the node and dilute the irrelevant parts to learn richer features, so we propose a disentangled knowledge graph aggregation module, which is divided into two parts : 1) Similarity-Aware Aggregation, 2) Mutual Information-Based Disentanglement.

Similarity-Aware Aggregation. In order to estimate the importance of each neighbor node in aggregation, we propose a similarity-aware aggregation strategy. It calculates the similarity of neighbor nodes on each potential spatial component. The

greater the similarity, the more likely there is a connection between the two nodes, so as to determine whether to aggregate the two nodes. The correlation between nodes and neighbor nodes is estimated by each potential spatial component:

$$h_{i,k}^{l+1} = \sigma\left(\sum_{(i,c) \in \hat{N}(u)} \alpha_{(i,j,c)}^k \phi(h_{j,k}^l, h_c^l, \theta_c)\right) \quad (2)$$

$$\begin{aligned} \alpha_{(i,j,c)}^k &= softmax((e_{i,c}^k)^T \cdot e_{j,c}^k) \\ &= \frac{exp((e_{i,c}^k)^T \cdot e_{j,c}^k)}{\sum_{(j',r) \in \hat{N}(i)} exp((e_{i,c}^k)^T \cdot e_{j',c}^k)} \end{aligned} \quad (3)$$

$$\phi(h_e, h_c, \theta_c) = (\theta_c \bullet h_e) - h_r \quad (4)$$

where $h_{i,k}^{l+1}$ denotes the representation of q_i on the k -th component after passing through the l -th layer, $\theta_c = diag(w_c)$ represents the projection matrix of the concept c , $e_{i,c}^k = h_{i,k} \circ \theta_c$ denotes the relation-aware weight of q_i on the k -th component, and $\hat{N}(i)$ denotes the neighbor node of q_i and itself. Similarly, h_c^l is the representation of concept c after passing through the l -th layer. h_c^l is updated with layer-wise linear transformation with parameter W_c^l :

$$h_c^{l+1} = W_c^l \bullet h_c^l \quad (5)$$

By calculating the similarity of neighbor nodes on K components and aggregating nodes with large similarity, the output question vector $\tilde{q} = \{h_{q,1}^{l+1}, h_{q,2}^{l+1}, \dots, h_{q,K}^{l+1}\}$ and the concept vector $\mathcal{E}_c^{ind} = h_c^{l+1}$ which is learned from the indirect relationship are obtained. However, there is not only a relationship between concepts and questions, but also a connection between concepts. For example, there may be a prerequisite relationship between concepts, where concept 1 is the prior knowledge of concept 2. In order to capture this relationship, we use a graph attention network:

$$\mathcal{E}_c^{dir} = \sum_{c' \in N_c} \alpha_{c,c'} W_{dir} \mathcal{E}_{c'} \quad (6)$$

$$\alpha_{c,c'} = Softmax_{c'}(att_{dir}([W_{dir} \mathcal{E}_c, W_{dir} \mathcal{E}_{c'}])) \quad (7)$$

where $\mathcal{E}_c$ is the input vector of concept c , $c' \in N_c$, N_c is the neighbor node set of concept c , W_{dir} is a trainable parameter, $\alpha_{c,c'}$ is the attention weight, $att \bullet$ is a linear layer with LeakyReLU activation function and $\bullet$ is a join operation.

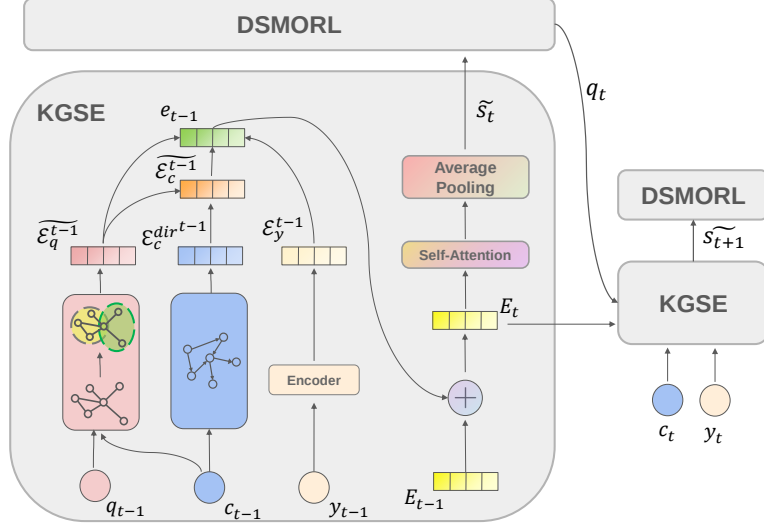


Figure 2: The structure of KGAT.

To distinguish the different information contained in $\mathcal{E}_c^{dir}$ and $\mathcal{E}_c^{ind}$, the attention vector $\mathbf{P}$ is used to calculate the weights of them respectively :

$$\begin{aligned} \mu_{ind} &= P_{ind}^T \bullet \tanh(W_{ind} \bullet \mathcal{E}_c^{ind} + b_{ind}) \\ \mu_{dir} &= P_{dir}^T \bullet \tanh(W_{dir} \bullet \mathcal{E}_c^{dir} + b_{dir}) \end{aligned} \quad (8)$$

Finally, the concept vector $\tilde{\mathcal{E}}_c$ can be defined as :

$$\tilde{\mathcal{E}}_c = \mu_{ind} \mathcal{E}_c^{ind} + \mu_{dir} \mathcal{E}_c^{dir} \quad (9)$$

Mutual Information-Based Disentanglement.

Although each question is mapped to K different concept components, there is still a weak correlation between these components and the relationship between these components is not only linear. Inspired by (Cheng et al., 2020b) and (Wu et al., 2021), we propose utilizing mutual information to achieve disentanglement. Mutual information can measure the degree of nonlinear dependence between two random variables. Here, the decoupling between components is achieved by using the contrastive log-ratio upper-bound MI estimator (Cheng et al., 2020a). Since the conditional probability of the same question on different conceptual components cannot be obtained directly, a simple neural network Q with variational distribution is proposed to approximate the real conditional neural network. The objective function is :

$$\begin{aligned} \mathcal{L}_{mi} &= \sum_u \sum_v E_{(h_{i,u}, h_{i,v}) \sim p(h_{i,u}, h_{i,v})} [\log q_\theta(h_{i,u} | h_{i,v})] \\ &\quad - E_{(h_{i,u}, h_{i',v}) \sim p(h_{i,u}, h_{i',v})} [\log q_\theta(h_{i,u} | h_{i',v})] \end{aligned} \quad (10)$$

Note that, $u \neq v$ and Q are trained simultaneously to minimize the KL divergence between $p(h_{i,u} | h_{i,v})$ and $q_\theta(h_{i,u} | h_{i,v})$.

$$\mathcal{L}_{(h_{i,u}, h_{i,v})} = \mathbb{D}_{KL}[p(h_{i,u} | h_{i,v}) || q_\theta(h_{i,u}, h_{i,v})] \quad (11)$$

Here, $p(h_{i,u} | h_{i,v})$ is a Gaussian distribution and alternately optimized with $\mathcal{L}_{mi}$. By using mutual information to propose the objective function, the dependence between different components is weakened, so that the relevant information is aggregated as much as possible when the adjacent nodes are aggregated, and the influence of irrelevant components is reduced.

State Encoder. We define a student's historical interaction sequence as $\tilde{s}_i = f(s_1^i, s_2^i, \dots, s_{i-1}^i)$, $s_{i-1}^i = (q_{i-1}^i, c_{i-1}^i, y_{i-1}^i)$ in reinforcement learning.

After the origin of the relation-aware question vector $\tilde{\mathcal{E}}_q$ and the concept vector $\tilde{\mathcal{E}}_c$ are obtained in Similarity-Aware Aggregation, the embedding vector $W_y = \mathbb{R}^{2 \times d}$ is used to transform the student's response x_y into a real-valued embedding $\mathcal{E}_y \in \mathbb{R}^d$:

$$\mathcal{E}_y = x_y W_y \quad (12)$$

where x_y is the one-hot encoding of response y . Therefore, $e_{t'}$ at time t' can be expressed as:

$$e_{t'} = \tilde{\mathcal{E}}_q^{t'} \oplus \tilde{\mathcal{E}}_c^{t'} \oplus \mathcal{E}_y^{t'} \quad (13)$$

where $t' \in [1, t-1]$, $e_{t'} \in \mathbb{R}^D$, $D = 3d$.

$$E_t = [e_1, e_2, \dots, e_{t-1}]^T \in \mathbb{R}^{(t-1) \times D} \quad (14)$$

Due to the varying amount of information contained in different responses, for example, answering a question correctly yields more information than answering a question incorrectly, a self attention mechanism is introduced here :

$$\widetilde{E}_t = \text{Softmax}\left(\frac{(E_t W^Q)(E_t W^K)^T}{\sqrt{d_k}}\right)(E_t W^V) \quad (15)$$

where W^Q, W^K, W^V are trainable parameters and $\sqrt{d_k}$ is the scale factor.

After the self-attention mechanism, LayerNorm and skip-connection are added, and Dropout is used to avoid overfitting. The self-attention results are processed using the average pooling layer to obtain the student state $\widetilde{s}_t \in \mathbb{R}^D$ for reinforcement learning.

3.3 DSMORL

Based on the multi-objective reinforcement learning framework considering accuracy, diversity and novelty proposed by (Wang et al., 2023), we add two other objectives: difficulty and similarity. Then Markov modeling based on multi-objective reinforcement learning is as follows :

$$\begin{aligned} \max_{\pi} \mathcal{J} &= \max_{\pi} \frac{1}{n} \sum_{i=1}^n [w^T (\sum_{t=1}^T \gamma^t r(s_t^i, q_t^i))] \\ &= \max_{\pi} E_{i \sim \pi} [w^T (\sum_{t=1}^T \gamma^t r(s_t^i, q_t^i))] \end{aligned} \quad (16)$$

where n is the number of students, q_t^i is the question selected according to the selection algorithm π in the state s_t^i , $r(s_t^i, q_t^i) = [r_{qua}, r_{div}, r_{nov}, r_{dif}, r_{sim}]$ is the reward obtained by selecting q_t^i and w is a weight vector used to calculate the importance of different reward components. Consideration factors are introduced by setting the reward values, including r_{qua}, r_{div} and r_{nov} the same as original MORL, in addition, the question difficulty r_{dif} and the student similarity r_{sim} are also introduced here.

Difficulty. It is generally believed that if students are given many very simple questions, their ability cannot be improved. If students are given too many very difficult questions, they may answer them casually because they are beyond their ability. Neither of these situations can reflect the true level of students. So the ability of student i as D_i :

$$D_i = \frac{1}{|Q_i|} \sum_{q \in Q_i} d_q, \quad d_q = \frac{1}{n} \sum_{j=1}^n r_{q,j} \quad (17)$$

where Q_i is the set of questions that student i has done, d_q is the difficulty of question q , the higher d_q is, the more difficult the question is, n is the number of times that question q is practiced by all students, and $r_{q,j}$ is the response of student j to question q . Assuming that the question q is selected at time t , the difficulty reward is defined as:

$$r_{dif} = \begin{cases} 1, & \text{if } |D_i - d_q| < x_{dif} \\ 0, & \text{otherwise} \end{cases} \quad (18)$$

where x_{dif} is a threshold set as 0.25 here.

Similarity. Using user similarity to recommend products for users is a commonly used method in recommendation systems. Specifically, if the shopping habits of user u_i and user u_j are similar, u_i is likely to be interested in what u_j has purchased, and the purchase probability is high. This idea can be used to recommend questions for students who may be interested in or that he may potentially need to practice. Here, the jaccard coefficient is used to measure the similarity between student s_i and student s_j :

$$\text{similarity}(s_i, s_j) = \frac{|Q_{s_i} \cap Q_{s_j}|}{|Q_{s_i} \cup Q_{s_j}|} \quad (19)$$

The candidate similar question set $\mathcal{Q}_i$ of s_i is obtained by all the question sets done by s_j whose similarity with s_i is higher than x_{sim} :

$$\begin{aligned} \mathcal{Q}_i &= (\bigcup_{j=1}^n Q_{s_j}) / Q_{s_i} \quad (i \neq j, \\ &\quad \text{similarity}(s_i, s_j) > x_{sim}) \end{aligned} \quad (20)$$

Assuming that the question q_t is selected at time t , the similarity reward is defined as:

$$r_{sim} = \begin{cases} 1, & \text{if } q_t \in \mathcal{Q}_i \\ 0, & \text{otherwise} \end{cases} \quad (21)$$

Then, a reward $r(s_t^i, q_t^i) = [r_{qua}, r_{div}, r_{nov}, r_{dif}, r_{sim}]$ is defined, and the weight of each reward component can be set by the binary vector w . We uses Actor-Critic as the recommender, policy network which is a fully connected layer with the parameter ϕ_{π} as the actor from the distribution $\pi(q_t|s_t; \phi_{\pi})$ sampling selection question and a fully connected layer value network with a parameter ϕ_v as a critic to evaluate the state. Given a state, the critic output vector $V(s_t; \phi_{\pi}) =$

$[V(s_t)_{qua}, V(s_t)_{div}, V(s_t)_{nov}, V(s_t)_{dif}, V(s_t)_{sim}]$ is used to predict the expected return. To maximize the weight sum $\mathcal{J}$ in Equation (16), multi-objective PPO (Schulman et al., 2017) is used here.

The advantage value of q_t is defined as the actual return of a state-action pair minus the predicted value of this state:

$$A(s_t, q_t) = \sum_{t'=t} \gamma^{t'-t} r(s_{t'}, q_{t'}) - V(s_t) \quad (22)$$

In order to convert A into a scalar, the weight vector w is used here, and each component in w corresponds to the quality, diversity, novelty, difficulty, and similarity objectives. Here, the clipped surrogate loss is applied to update the action parameters:

$$\mathcal{L}_1 = -\mathbb{E}_{\pi_{old}} [\text{Min}\{\frac{\pi(q_t|s_t)}{\pi_{old}(q_t|s_t)} w^T A(s_t, q_t), \text{Clip}(\frac{\pi(q_t|s_t)}{\pi_{old}(q_t|s_t)}, 1 - \epsilon, 1 + \epsilon) w^T A(s_t, q_t)\}] \quad (23)$$

The critic loss is proposed based on the expectation that the expected return is as close as possible to the actual return:

$$\mathcal{L}_2 = \frac{1}{2} w^T \|V(s_t) - \gamma^{t'-t} r(s_{t'}, q_{t'})\|^2 \quad (24)$$

The final multi-objective PPO loss function is :

$$\mathcal{L} = \mathcal{L}_1 + \alpha \mathcal{L}_2 \quad (25)$$

where α is the balance parameter.

4 Experiment

4.1 Dataset and Setting

We conduct experiments on three public datasets (Assist2009 (Feng et al., 2009), Junyi (Chang et al., 2015), Eedi (Wang et al., 2020b)). These datasets record the data of students' questions on the online education platform. After removing the data with the length of the student interaction sequence less than 40, the statistical data of the three datasets after processing are shown in Table I.

The experiment is divided into training, test and verification set according to the proportion of 80 % -10 % -10 %. At the same time, the data of each student is divided into candidate data set and meta-data set according to the proportion of 80 % -20 %. The data in the metadata set is considered to be the student's historical interactive records, and the data in the candidate data set may be recommended to the students according to the selection

DataSet	Assist2009	Junyi	Eedi
Student	1,360	20,395	4,918
Question	17,751	2,835	948
Concept	123	40	86
Interactions	239,919	2,537,898	1,382,727
Concepts per Question	1.2	1.0	4.0
Question per Concept	172.7	70.9	44.28
Positive Label Rate	0.55	0.62	0.69

Table 1: Statistics of the Datasets Used in the Experiments.

algorithm. These two data sets are randomly generated in each training phase to prevent overfitting. In the experiment, the weight factor $w = [1, 1, 1, 1, 1]$, the potential space projection number $K = 3$ and the results when the step=10, 20 are shown. Set $x_{dif} = 0.25$, $x_{sim} = 0.1$, $\gamma = 0.5$ in Equation (22), $\epsilon = 0.2$ in Equation (23), batch size is 128, embedding dimension $d = 16$. The dropout factor in the self-attention mechanism is 0.1. The optimizer is Adam optimizer, the learning rate is 0.02, and the loss balance factor is 1. The experiment directly uses the parameters of baselines in the original text to ensure their best effect. Since only Junyi provides the premise relationship, (Gao et al., 2021) is used to construct the concept directed graphs in the other two datasets.

4.2 Performance Comparison

To verify the effectiveness of KGCAT, the experiment compared KGCAT with MAAT (Bi et al., 2020), BOBCAT (Ghosh and Lan, 2021), NCAT (Zhuang et al., 2022) and GMOCAT (Wang et al., 2023), basing IRT and NCD respectively. In this work, the models predict the results of students' responses to questions (correct or incorrect). Table II and Table III are the results at step = 10 and 20 with AUC (area under the ROC curve) and ACC (accuracy) as metrics. The overall performance of KGCAT is better than baselines. This indicates that using knowledge graph to learn the semantic and relational information of questions and concepts, as well as considering the difficulty of questions and the similarity of students, is beneficial for improving the performance of CAT. KGCAT sometimes performs worse than baselines with fewer steps, but as the steps increase, its AUC and ACC will continue to increase and reach their optimal. This proves the superiority of reinforcement learning in longer-term scenarios. At the same time, NCAT, GMOCAT, KGCAT, which use reinforcement learning methods, are superior to other meth-

Dataset	Assist2009				Junyi				Eedi			
CDM	IRT		NCD		IRT		NCD		IRT		NCD	
step	10	20	10	20	10	20	10	20	10	20	10	20
MAAT	68.82	69.70	69.38	71.17	77.39	78.21	77.53	78.40	70.90	73.19	71.03	73.75
BOBCAT	69.44	70.97	70.63	71.80	78.70	79.17	78.21	79.46	70.50	73.24	71.44	74.51
NCAT	69.49	71.06	70.93	71.68	78.87	79.13	78.41	79.56	70.78	73.32	71.45	74.55
GMOCAT	70.02	72.08	71.33	72.99	79.44	80.19	79.60	80.09	73.23	75.30	73.57	75.49
KGCAT	71.15	72.50	71.98	73.62	79.85	80.94	79.84	80.95	73.58	75.68	73.63	75.91

Table 2: AUC comparison results of all the baselines and proposed models on 3 Datasets. The best results are in bold and the second best results are underlined.

Dataset	Assist2009				Junyi				Eedi			
CDM	IRT		NCD		IRT		NCD		IRT		NCD	
step	10	20	10	20	10	20	10	20	10	20	10	20
MAAT	66.30	67.57	67.84	69.52	74.51	75.48	74.88	75.38	64.38	66.15	64.58	66.71
BOBCAT	66.21	67.88	68.63	69.94	<u>75.66</u>	<u>76.51</u>	75.21	76.45	64.90	66.97	65.62	67.69
NCAT	66.32	68.36	68.38	69.44	76.05	75.70	75.47	76.40	64.97	66.92	65.59	67.84
GMOCAT	<u>67.32</u>	<u>68.99</u>	<u>69.31</u>	<u>70.63</u>	75.07	76.23	<u>75.60</u>	<u>76.72</u>	<u>66.81</u>	<u>68.55</u>	<u>67.06</u>	<u>69.01</u>
KGCAT	67.81	69.67	69.75	71.28	75.51	76.90	75.61	76.98	67.00	68.85	67.20	69.18

Table 3: ACC comparison results of all the baselines and proposed models on 3 Datasets. The best results are in bold and the second best results are underlined.

ods that do not use reinforcement learning, which can also prove the effectiveness of reinforcement learning in the application of selection algorithms. The improvement of KGCAT on Assist2009 and Junyi is higher than that on Eedi. This is because on the Assist2009 and Junyi datasets, each question involves an average of 1.2 and 1 concepts. When we project each question to 3 latent spaces, more conceptual information is captured, while on Eedi, each question involves an average of 4 concepts, and projecting the question to 3 latent spaces may cause information loss.

4.3 Ablations and Discussion

In order to prove the effectiveness of each module of KGCAT, the ablation experiment was carried out on the Eedi dataset. KGCAT-KG, KGCAT-H, and KGCAT-S represent the models of KGCAT replacing knowledge graph with ordinary GNN, removing question difficulty and student similarity rewards respectively. By setting $w = [1, 1, 1, 0, 1]$ and $w = [1, 1, 1, 1, 0]$ in Equation (16), KGCAT-H and KGCAT-S were realized respectively. In Figure 3, the performance of KGCAT decreased regardless of which module was removed when step = 20, which proved that each module plays a role in improving the performance of the model. Among them, the performance of KGCAT-H and KGCAT-S decreased rapidly, indicating the impor-

tance of question difficulty and student similarity in the selection of questions, and further proving the superiority of reinforcement learning. In order to explore the optimal values of x_{dif} and s_{sim} , we conducted additional experiments on Eedi dataset.

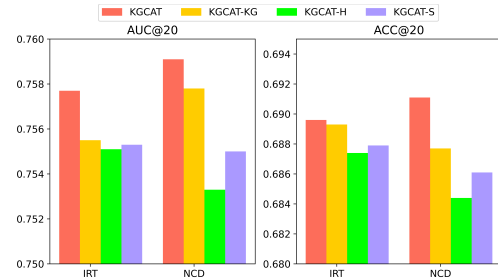


Figure 3: The Results of Ablation Experiments.

Projection number K analysis. In Section 3.2, it is mentioned that in order to learn more abundant features, the questions are projected into K potential spaces. In order to explore the most suitable K value, we conducted experiments. Due to the limitation of computing power, K here only takes values from 1 to 3. The results are shown in Figure 4. When K = 3, the results are the best, which is in line with cognition. The more space the questions are projected into, the more different semantics can be learned and the performance of the model can be improved. The influence of KGCAT-KG on the model in the ablation experiment is not as good

as that of KGCAT-H and KGCAT-S. It may also because the value of K is not large enough to learn richer semantic information.

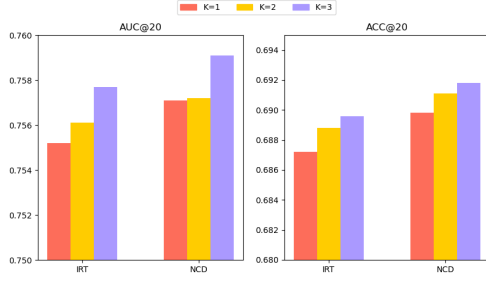


Figure 4: AUC and ACC Performance Comparison with Different K value.

x_{dif} analysis. Figure 5 shows the distribution of question difficulty on Eedi. Most of the questions are of medium difficulty, which is consistent with the distribution of question difficulty in most education platforms. The experimental results are shown in Figure 6. AUC and ACC have the best effect when $x_{dif} = 0.25$, and the effect gradually deteriorates as the values are taken at both ends. This is because if the value of x_{dif} is smaller, the fewer questions meeting the conditions, resulting in part of the questions that are actually useful for the model are also filtered, making the effect worse. However, if the value of x_{dif} is too large, which will reduce the screening effect on the questions, which will also reduce the performance of the model.

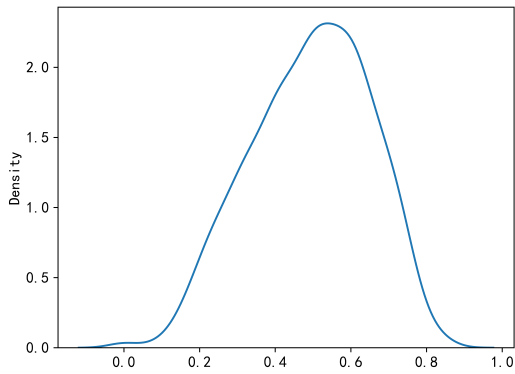


Figure 5: The Probability Density Distribution of Question Difficulty.

x_{sim}	0.1	0.2	0.3	0.4	0.5
Students	4,914	4,908	4,875	4,766	4,463

Table 4: The Number of Students with Different x_{sim} .

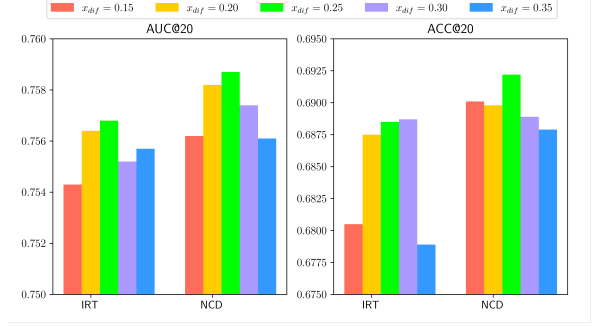


Figure 6: AUC and ACC Performance Comparison with Different x_{dif} .

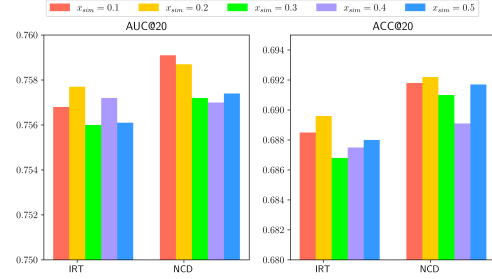


Figure 7: AUC and ACC Performance Comparison with Different x_{sim} .

x_{sim} analysis. Table IV is the number of students whose length of candidate similar question sets Q_i is greater than 0 under different x_{sim} . When $x_{sim} > 0.4$, the number of students with candidate similar question sets begins to decline rapidly. In order to make as many students as possible have their own candidate similar question sets, the AUC and ACC of the model were calculated by using $x_{sim} \leq 0.5$. As shown in Figure 7, when $x_{sim} = 0.1/0.2$, the overall performance is better. Meanwhile, the value of x_{sim} should not be too large, and as many students as possible should have their own candidate similar question sets.

5 Conclusion

To enhance the effectiveness of personalized education recommendations, we propose a Knowledge Graph Reasoning-Based Model for Computerized Adaptive Testing (KGCAT). This model leverages the reasoning ability of the knowledge graph, so as to deeply learn the relationship between questions and concepts and introduces a disentanglement module that uses mutual information to reduce noise by minimizing the influence of low-relevance neighbor nodes. Additionally, by incorporating reinforcement learning into the decision-making process, the model considers both question difficulty

and student similarity in a multi-objective manner, allowing for more accurate question recommendations. Finally, a series of experiments demonstrate the model’s effectiveness and improved interpretability

6 Limitation

While KGCAT effectively utilizes mutual information to reduce noise and improve recommendation accuracy, processing large-scale knowledge graphs with numerous concepts and questions could introduce computational challenges, potentially limiting scalability. Additionally, the integration of mutual information for concept disentanglement and multi-objective reinforcement learning introduces added complexity to the hyperparameter tuning process, requiring considerable time and effort to optimize components such as learning rates, reward structures, and mutual information thresholds.

References

- Qingyao Ai, Vahid Azizi, Xu Chen, and Yongfeng Zhang. 2018. Learning heterogeneous knowledge base embeddings for explainable recommendation. *Algorithms*, 11(9):137.
- Haoyang Bi, Haiping Ma, Zhenya Huang, Yu Yin, Qi Liu, Enhong Chen, Yu Su, and Shijin Wang. 2020. Quality meets diversity: A model-agnostic framework for computerized adaptive testing. In *2020 IEEE International Conference on Data Mining (ICDM)*, pages 42–51. IEEE.
- Yixin Cao, Xiang Wang, Xiangnan He, Zikun Hu, and Tat-Seng Chua. 2019. Unifying knowledge graph learning and recommendation: Towards a better understanding of user preferences. In *The world wide web conference*, pages 151–161.
- Rose Catherine and William Cohen. 2016. Personalized recommendations using knowledge graphs: A probabilistic logic programming approach. In *Proceedings of the 10th ACM conference on recommender systems*, pages 325–332.
- Haw-Shiuan Chang, Hwai-Jung Hsu, Kuan-Ta Chen, et al. 2015. Modeling exercise relationships in e-learning: A unified approach. In *EDM*, pages 532–535.
- Suming Jeremiah Chen, Arthur Choi, and Adnan Darwiche. 2015. Computer adaptive testing using the same-decision probability. In *BMA@ UAI*, pages 34–43.
- Pengyu Cheng, Weituo Hao, Shuyang Dai, Jiachang Liu, Zhe Gan, and Lawrence Carin. 2020a. Club: A contrastive log-ratio upper bound of mutual information. In *International conference on machine learning*, pages 1779–1788. PMLR.
- Pengyu Cheng, Martin Renqiang Min, Dinghan Shen, Christopher Malon, Yizhe Zhang, Yitong Li, and Lawrence Carin. 2020b. Improving disentangled text representation learning with information-theoretic guidance. *arXiv preprint arXiv:2006.00693*.
- Susan E Embretson and Steven P Reise. 2013. *Item response theory*. Psychology Press.
- Mingyu Feng, Neil Heffernan, and Kenneth Koedinger. 2009. Addressing the assessment challenge with an online system that tutors as it assesses. *User modeling and user-adapted interaction*, 19:243–266.
- Lina Gao, Zhongying Zhao, Chao Li, Jianli Zhao, and Qingtian Zeng. 2022. Deep cognitive diagnosis model for predicting students’ performance. *Future Generation Computer Systems*, 126:252–262.
- Weibo Gao, Qi Liu, Zhenya Huang, Yu Yin, Haoyang Bi, Mu-Chun Wang, Jianhui Ma, Shijin Wang, and Yu Su. 2021. Rcd: Relation map driven cognitive diagnosis for intelligent education systems. In *Proceedings of the 44th international ACM SIGIR conference on research and development in information retrieval*, pages 501–510.
- Aritra Ghosh and Andrew Lan. 2021. Bobcat: Bilevel optimization-based computerized adaptive testing. *arXiv preprint arXiv:2108.07386*.
- Hiromi Nakagawa, Yusuke Iwasawa, and Yutaka Matsuo. 2019. Graph-based knowledge tracing: modeling student proficiency using graph neural network. In *IEEE/WIC/ACM International Conference on Web Intelligence*, pages 156–163.
- Franco Scarselli, Marco Gori, Ah Chung Tsoi, Markus Hagenbuchner, and Gabriele Monfardini. 2009. [The graph neural network model](#). *IEEE Transactions on Neural Networks*, 20:61–80.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. 2017. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*.
- Hanshuang Tong, Zhen Wang, Yun Zhou, Shiwei Tong, Wenyuan Han, and Qi Liu. 2022. Introducing problem schema with hierarchical exercise graph for knowledge tracing. In *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pages 405–415.
- Fei Wang, Qi Liu, Enhong Chen, Zhenya Huang, Yuying Chen, Yu Yin, Zai Huang, and Shijin Wang. 2020a. Neural cognitive diagnosis for intelligent education systems. In *Proceedings of the AAAI conference on artificial intelligence*, volume 34, pages 6153–6161.

- Hangyu Wang, Ting Long, Liang Yin, Weinan Zhang, Wei Xia, Qichen Hong, Dingyin Xia, Ruiming Tang, and Yong Yu. 2023. Gmocat: A graph-enhanced multi-objective method for computerized adaptive testing. In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pages 2279–2289.
- Xiang Wang, Xiangnan He, Yixin Cao, Meng Liu, and Tat-Seng Chua. 2019a. Kgat: Knowledge graph attention network for recommendation. In *Proceedings of the 25th ACM SIGKDD international conference on knowledge discovery & data mining*, pages 950–958.
- Xiang Wang, Dingxian Wang, Canran Xu, Xiangnan He, Yixin Cao, and Tat-Seng Chua. 2019b. Explainable reasoning over knowledge graphs for recommendation. In *Proceedings of the AAAI conference on artificial intelligence*, volume 33, pages 5329–5336.
- Zichao Wang, Angus Lamb, Evgeny Saveliev, Pashmina Cameron, Yordan Zaykov, José Miguel Hernández-Lobato, Richard E Turner, Richard G Baraniuk, Craig Barton, Simon Peyton Jones, et al. 2020b. Instructions and guide for diagnostic questions: The neurips 2020 education challenge. *arXiv preprint arXiv:2007.12061*.
- David J Weiss and G Gage Kingsbury. 1984. Application of computerized adaptive testing to educational problems. *Journal of educational measurement*, 21(4):361–375.
- Junkang Wu, Wentao Shi, Xuezhi Cao, Jiawei Chen, Wenqiang Lei, Fuzheng Zhang, Wei Wu, and Xiangnan He. 2021. [Disenkgat: Knowledge graph embedding with disentangled graph attention network](#). *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*.
- Yang Yang, Jian Shen, Yanru Qu, Yunfei Liu, Kerong Wang, Yaoming Zhu, Weinan Zhang, and Yong Yu. 2021. Gikt: a graph-based interaction model for knowledge tracing. In *Machine Learning and Knowledge Discovery in Databases: European Conference, ECML PKDD 2020, Ghent, Belgium, September 14–18, 2020, Proceedings, Part I*, pages 299–315. Springer.
- Yan Zhuang, Qi Liu, Zhenya Huang, Zhi Li, Shuanghong Shen, and Haiping Ma. 2022. Fully adaptive framework: Neural computerized adaptive testing for online education. In *Proceedings of the AAAI conference on artificial intelligence*, volume 36, pages 4734–4742.