

Prompt Transfer for Dual-Aspect Cross-Domain Cognitive Diagnosis

Fei Liu , Yizhong Zhang , Shuochen Liu , Shengwei Ji , Kui Yu , and Le Wu 

Abstract—Cognitive diagnosis (CD) aims to evaluate students' cognitive states based on their interaction data, enabling downstream applications such as exercise recommendation and personalized learning guidance. However, existing methods often struggle with accuracy drops in cross-domain cognitive diagnosis (CDCD), a practical yet challenging task. While some efforts have explored exercise-aspect CDCD, such as cross-subject scenarios, they fail to address the broader dual-aspect nature of CDCD, encompassing both student- and exercise-aspect variations. This diversity creates significant challenges in developing a scenario-agnostic framework. To address these gaps, we propose PromptCD, a simple yet effective framework that leverages soft prompt transfer for cognitive diagnosis. PromptCD is designed to adapt seamlessly across diverse CDCD scenarios, introducing PromptCD-S for student-aspect CDCD and PromptCD-E for exercise-aspect CDCD. Extensive experiments on real-world datasets demonstrate the robustness and effectiveness of PromptCD, consistently achieving superior performance across various CDCD scenarios. Our work offers a unified and generalizable approach to CDCD, advancing both theoretical and practical understanding in this critical domain. The implementation of our framework is publicly available at <https://github.com/Publisher-PromptCD/PromptCD>.

Index Terms—Cognitive diagnosis, cross-domain, educational data mining, prompt transfer.

I. INTRODUCTION

COGNITIVE diagnosis aims to assess students' proficiency based on historical interactions [1], [2], [3]. It is a crucial task in educational data mining, supporting many downstream tasks such as exercise recommendation [4], [5], learning guidance [6], [7], [8], [9], [10], and computerized adaptive testing [11].

In recent efforts, the primary focus has been on enhancing the accuracy of cognitive diagnosis models [12]. Specifically, these models aim to learn the characteristics of students and exercises from training data and utilize these learned representations to predict scores on test data [13], [14], [15], [16], [17], [18]. Despite advancements in these models, they rely on the assumption that student and exercise characteristics are consistent across training and test data, which can be referred to as in-domain cognitive diagnosis. However, the assumption mentioned above is strict in practice. Students and exercises with diverse characteristics create diverse domains. For instance, students from different schools or countries, while exercises span various subjects. Thus, considering cross-domain cognitive diagnosis (CDCD) for students or exercises with various characteristics is more practical, however, posing many challenges to existing models.

First, representations of students or exercises learned from training data (source domains) cannot be directly applied to testing data (target domains), leading to a sharp decline in diagnostic accuracy. A straightforward approach to this limitation is to retrain the model. However, retraining solely on new domain data often results in overfitting and catastrophic forgetting, severely compromising the model's generalization capabilities. Alternatively, retraining with all existing and new data is computationally intensive and impractical for real-world applications. To illustrate these limitations, we conducted experiments using MIRT [19] on the SLP dataset [20] in in-domain (A and B) and cross-domain (C and D) scenarios. Fig. 1 provides detailed scenario descriptions, and the results highlight the following: 1) *Cross-Domain Challenges*: C and D performed significantly worse than A and B, indicating the struggles of traditional models in CDCD scenarios; 2) *Overfitting Risks*: A performed worse than B, highlighting the impact of overfitting when retraining with limited data; and 3) *Sensitivity to Domain Differences*: D significantly underperformed C due to greater distributional differences between Chinese and mathematics compared with physics and mathematics, underscoring the sensitivity of traditional models to source domains.

Received 30 November 2024; revised 2 April 2025; accepted 3 June 2025. Date of publication 8 July 2025; date of current version 3 December 2025. This work was supported in part by the National Science and Technology Major Project under Grant 2021ZD0111802; in part by the National Natural Science Foundation of China under Grant 62406096, Grant 72188101, Grant 62376086, Grant U23B2031, Grant 721881011, and Grant 62306100; in part by the China Postdoctoral Science Foundation under Grant 2024M760722; in part by Anhui Postdoctoral Scientific Research Program Foundation under Grant 2024C934; in part by the Key Laboratory of Knowledge Engineering with Big Data (the Ministry of Education of China) under Grant BigKEOpen2025-01; and in part by the Fundamental Research Funds for the Central Universities under Grant JZ2024HGQB0093. (Corresponding author: Le Wu.)

Fei Liu, Yizhong Zhang, Kui Yu, and Le Wu are with the Key Laboratory of Knowledge Engineering with Big Data (the Ministry of Education of China), Hefei University of Technology, Hefei 230601, P.R. China, and also with the School of Computer Science and Information Engineering, Hefei University of Technology, Hefei 230601, P.R. China (e-mail: feiliu@mail.hfut.edu.cn).

Shuochen Liu is with the School of Computer Science and Technology, University of Science and Technology of China, Hefei 230026, China.

Shengwei Ji is with the Key Laboratory of Knowledge Engineering with Big Data (the Ministry of Education of China), Hefei University of Technology, Hefei 230601, P.R. China, and also with the School of Big Data and Artificial Intelligence, Hefei University, Hefei 230601, China.

Digital Object Identifier 10.1109/TCSS.2025.3577255

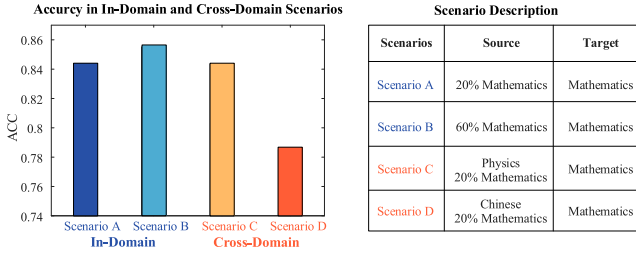


Fig. 1. Performance comparison of MIRT in in-domain scenarios A and B versus cross-domain scenarios C and D.

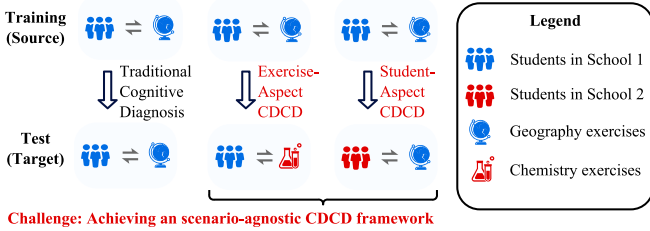


Fig. 2. Traditional cognitive diagnosis versus cross-domain cognitive diagnosis.

These observations demonstrate that retraining cognitive diagnosis models is impractical, necessitating a better approach to mitigate accuracy declines in CDCD.

Second, CDCD scenarios are inherently complex due to the diverse characteristics of students and exercises, creating a variety of domains. As shown in Fig. 2, CDCD can be categorized into two types: 1) *Student-Aspect CDCD*: Differences arise from varying student demographics, such as urban versus rural populations; and 2) *Exercise-Aspect CDCD*: Variations occur across subject domains, such as mathematics and physics. Unfortunately, CDCD remains an underexplored area. Existing studies [21], [22] have primarily focused on specific CDCD aspects, proposing scenario-specific models with limited compatibility. Developing a generalizable framework capable of addressing both student- and exercise-aspect CDCD scenarios remains a significant challenge.

In this article, we propose PromptCD, a simple yet generalizable framework designed to address the challenges of dual-aspect cross-domain cognitive diagnosis (CDCD). Dual-aspect CDCD introduces unique complexities, as knowledge transfer must consider both student-aspect and exercise-aspect scenarios. These scenarios involve distinct challenges. 1) Students and exercises across domains can be either overlapping or nonoverlapping. Overlapping entities require personalized adaptation to maintain consistency, while nonoverlapping entities necessitate a generalized representation to ensure effective knowledge transfer. 2) Diverse target domains vary significantly in their characteristics, making it difficult to adapt representations while preserving diagnostic accuracy and avoiding issues such as overfitting or catastrophic forgetting. To address these challenges, PromptCD introduces a unified framework leveraging soft prompt transfer, a proven technique in

cross-domain tasks across various fields [23], [24], [25], [26], [27], [28]. Specifically, we design personalized prompts for overlapping entities and shared domain-adaptive prompts for nonoverlapping entities. These prompts enhance representation learning and transfer, ensuring robustness in diverse CDCD scenarios. Additionally, PromptCD adopts a two-stage training strategy—pretraining on source domains and fine-tuning on target domains—for efficient and scalable adaptation. To demonstrate its versatility, we develop PromptCD-S and PromptCD-E, tailored to student-aspect and exercise-aspect CDCD scenarios, respectively. We summarize the contributions of this article as follows.

- 1) We propose the PromptCD framework, introducing soft prompt transfer and a two-stage training strategy to address dual-aspect CDCD challenges.
- 2) We develop PromptCD-S and PromptCD-E, showcasing the framework's ability to generalize across student- and exercise-aspect scenarios.
- 3) Extensive experiments on real-world datasets validate the effectiveness of PromptCD, achieving significant performance improvements over baselines.

II. PRELIMINARIES

A. Cognitive Diagnosis

Cognitive diagnosis aims to evaluate students' proficiency in knowledge concepts based on their response records L . This task involves modeling the interaction between student features α and exercise features β to predict scores. Since students' proficiency is not directly observable, the model is trained to optimize predictive accuracy using the cross-entropy loss $\mathcal{L}_{CE}$. Below, we outline key interaction functions used in classic cognitive diagnosis models.

IRT [29] and MIRT [19] use the logistic function in a unidimensional and multidimensional manner, respectively. The detailed interaction functions are as follows: $y_{uv} = 1/(1 + e^{-C \cdot D_v(\alpha_u - \beta_v)})$ and $y_{uv} = 1/(1 + e^{-\alpha_u^T \beta_v + D_v})$, where D_v is discrimination of exercise v . C is a constant. NeuralCD [2], [15] utilizes neural networks to model the complex interactions between representations of students and exercises as follows: $x_{uv} = Q_v \circ (\alpha_u - \beta_v) * D_v, y_{uv} = f_1(f_2(f_3(x_{uv})))$, where $Q_v \in \{0, 1\}^{1 \times K}$ indicates whether an exercise is associated with a knowledge concept. f_1, f_2, f_3 are the fully connected layers with positive weights to ensure monotonicity. KSCD [3] further explores the impact of potential associations between knowledge concepts on diagnostic results, shown as follows: $\alpha'_{uc} = \phi(f_{sk}(\alpha_u \oplus h_c^K)), \beta'_{vc} = \phi(f_{ek}(\beta_v \oplus h_c^K)), y_{uv} = \phi(\frac{1}{n_v} \sum_{c=1}^C Q_{vc} \times f_{se}(\alpha'_{uc} - \beta'_{vc}))$, where h_c^K represents the initialized embedding representations of knowledge concept c . Q_{vc} is the knowledge relevance vector Q_v of the concept c . n_v indicates the number of knowledge concepts contained in exercise e_v . ϕ is the activation function. f_{se}, f_{sk}, f_{ek} are linear transformation functions that correspond to different fully connected layers.

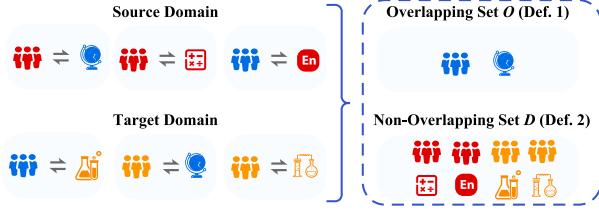


Fig. 3. Illustration of overlapping and non-overlapping sets in CDCD.

B. Cross-Domain Cognitive Diagnosis (CDCD)

Consider $|S|$ source domains $\{S_1, S_2, \dots, S_{|S|}\}$ and $|T|$ target domains $\{T_1, T_2, \dots, T_{|T|}\}$. Let LS_s and LT_t denote the interaction records for source domain S_s and target domain T_t , respectively, where $s \in \{1, 2, \dots, |S|\}$ and $t \in \{1, 2, \dots, |T|\}$. The CDCD task aims to identify and leverage cognitive patterns and learning structures that generalize across domains, enhancing the model's performance in the new domain T_t . By leveraging the abundant data $LS_1, LS_2, \dots, LS_{|S|}$ from source domains, CDCD can rapidly establish cognitive diagnosis models in target domain T_t with few-shot data $LT_t^{few} \subset LT_t$.

To facilitate the introduction of the subsequent framework, we define sets O and D to represent the overlapping and nonoverlapping entities between the source and target domains, as depicted in Fig. 3. The entity here refers to students or exercises. For instance, in the exercise-aspect CDCD scenarios, there exist nonoverlapping groups of exercises, defined as D . Conversely, we define the set of overlapping students as O , which allows the transfer of cross-domain information.

Definition 1 (Overlapping Set): The overlapping set O represents the entities that exist in both the source and target domains, which is defined as

$$O = \bigcup_{s=1}^{|S|} S_s \cap \bigcup_{t=1}^{|T|} T_t. \quad (1)$$

Definition 2 (Nonoverlapping Set): Let Ω be the universal set of all entities in both the source and target domains. The nonoverlapping set D is defined as the complement of O with respect to the universal set Ω

$$D = \Omega \setminus O = \left(\bigcup_{s=1}^{|S|} S_s \cup \bigcup_{t=1}^{|T|} T_t \right) \setminus O. \quad (2)$$

III. PROPOSED FRAMEWORK

In this section, we propose the scenario-agnostic PromptCD, applicable to both student- and exercise-aspect CDCD scenarios.

A. Overall Architecture

The overall architecture is shown in Fig. 4. Our two-stage framework abstracts scenario-agnostic features and unified learning strategies, enabling rapid adaptation to new domains. The pseudocode is presented in Algorithm 1.

Algorithm 1: PromptCD.

```

1: Input: cognitive diagnosis model  $\mathcal{M}$ , records  $LS$  for pre-
   training and  $LT_t^{few}$  in target domain  $t$  for fine-tuning.
2: Output: fine-tuned model  $\mathcal{M}$ , the transfer prompts  $\hat{p}^o$  and  $\hat{p}_t^d$ .
3: —Pre-training Stage—
4: while  $e_1 \leq Epoch_{Pretrain}$  do
5:   for  $LS_s \in \{LS_1, LS_2, \dots, LS_{|S|}\}$  do
6:     Initialize embeddings  $\mathbf{o}_s^{orig}$ ,  $\mathbf{d}_s^{orig}$ , prompts  $\mathbf{p}^o$ ,  $\mathbf{p}^d$  and  $\mathcal{M}$ ;
7:     Enhance the representation of entities in Eq.(3) and Eq.(4);

8:     Input  $\mathbf{o}_s^{out}$  and  $\mathbf{d}_s^{out}$  to  $\mathcal{M}$  to predict scores  $\mathbf{y}_s$ ;
9:     Calculate the loss using  $LS_s$  to update the model;
10:   end for
11: end while
12: —Fine-Tuning Stage—
13: while  $e_2 \leq Epoch_{Finetune}$  do
14:   Initialize the entities  $\mathbf{o}_t^{orig}$ ,  $\mathbf{d}_t^{orig}$ ;
15:   Obtain the transfer prompts  $\hat{p}^o$  and  $\hat{p}_t^d$  in Eq.(5) and Eq.(6);
16:   Activate improvement policy in Eq.(7);
17:   Enhance the representations in a manner similar to pre-
   training;
18:   Input  $\mathbf{o}_t^{out}$  and  $\mathbf{d}_t^{out}$  to  $\mathcal{M}$  to predict scores  $\mathbf{y}_t$ ;
19:   Calculate the loss using  $LT_t^{few}$  to update the model.
20: end while

```

Pretraining Stage: In Section III-B, we present an exposition of the personalized and shared prompts, as well as the processing strategies for entity representations. During the pretraining stage, the prompts are updated using the data LS_s ($s \in \{1, 2, \dots, |S|\}$) from the source domains.

Fine-Tuning Stage: In Section III-C, we outline the prompt transfer process and introduce a variant that enhances adaptation. After pretraining, we fine-tune the trainable parameters using few-shot data LS_t^{few} from the target domain T_t , thereby transferring knowledge from the source domains and adapting to the target domain's distribution.

B. Source Prompt Enhancement

Considering the characteristics of overlapping and nonoverlapping entities in cross-domain scenarios, relying solely on entity representations may not accommodate the diverse interactions across different domains. Consequently, we designed two types of learnable soft prompts—personalized prompts and shared prompts—to establish connection between the source and target domains by associating these prompts with different entity representations.

Specifically, personalized prompts are tailored for individual entities within the overlapping set O . Since the overlapping entities are identical in both the source and target domains, each entity can be associated with a personalized prompt, allowing for the transfer of more information. In contrast, shared prompts are utilized by all entities in each domain within the nonoverlapping set D , ensuring a common representation for domain-specific knowledge. The resulting composite representations can be expressed as

$$\mathbf{o}_{k,i}^{cat} = [\mathbf{p}_i^o, \mathbf{o}_{k,i}^{orig}], \quad \mathbf{d}_{k,j}^{cat} = [\mathbf{p}_k^d, \mathbf{d}_{k,j}^{orig}] \quad (3)$$

where $\mathbf{p}_i^o$ is the personalized prompt for each individual overlapping entity i , and $\mathbf{p}_k^d$ is the shared prompt for domain-specific

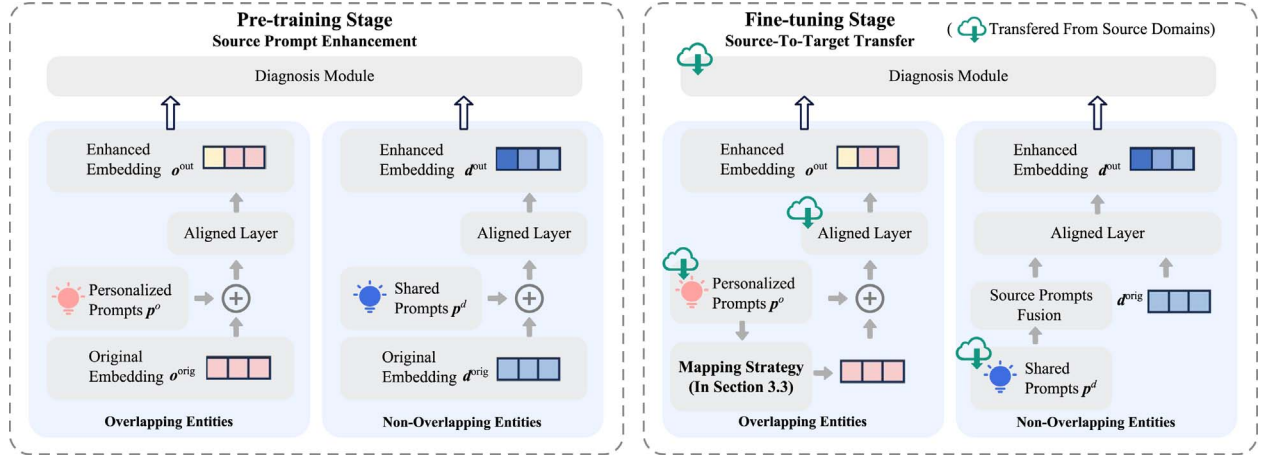


Fig. 4. Overall architecture of the proposed PromptCD framework, including the pretraining and fine-tuning stages.

entities. $\mathbf{o}_{k,i}^{\text{orig}}$ and $\mathbf{d}_{k,j}^{\text{orig}}$ respectively denote original embedding of single entity in source domain $\mathcal{S}_k$ ($k \in 1, 2, \dots, |\mathcal{S}|$) through random initialization, while $\mathbf{o}_{k,i}^{\text{cat}}$ and $\mathbf{d}_{k,i}^{\text{cat}}$ denote their corresponding representations after concatenation with prompts.

To integrate the concatenated features and extract the joint information, we utilize a fully connected layer defined as the operator linear, which has different trainable parameters depending on the types of entities

$$\mathbf{o}_k^{\text{out}} = \text{Linear}_o(\mathbf{o}_k^{\text{cat}}), \quad \mathbf{d}_k^{\text{out}} = \text{Linear}_d(\mathbf{d}_k^{\text{cat}}). \quad (4)$$

$\mathbf{o}_k^{\text{out}}$ and $\mathbf{d}_k^{\text{out}}$ denote final source-domain representations, aligned with the original embeddings. The processed representations are then input into the cognitive diagnosis model to predict scores.

C. Source-to-Target Transfer

In the fine-tuning stage, learned prompts are adapted to the target domains. Personalized prompts, designed for overlapping entities across domains, are transferred on a one-to-one basis in (5) to maintain the integrity of cross-domain connection information, as these prompts are also relevant to the target domains

$$\hat{\mathbf{p}}_i^o = \mathbf{p}_i^o. \quad (5)$$

To effectively capture the commonalities across different source domains, the shared prompts should be concatenated and mapped into a unified representation for the target domain. Specifically, we apply a learnable linear transformation to project the concatenated prompts into a shared latent space, ensuring both domain adaptation and dimension consistency. Formally, this process is expressed as

$$\hat{\mathbf{p}}_t^d = \text{Linear}_{s2t}(\mathbf{p}_1^d \oplus \mathbf{p}_2^d \oplus \dots \oplus \mathbf{p}_{|\mathcal{S}|}^d) \quad (6)$$

where Linear_{s2t} is a trainable mapping function. More specifically, we define it as $\hat{\mathbf{P}}_t^d = \mathbf{W}\mathbf{P}_S + \mathbf{b}$, where $\mathbf{P}_S$ is the

concatenated source-domain prompt matrix, $\mathbf{W}$ is the transformation matrix, and $\mathbf{b}$ is a bias term. This mapping serves to align feature distributions across domains while preserving essential shared knowledge, allowing the model to adaptively reweight different source-domain prompts for effective transfer.

The final representations $\mathbf{o}_t^{\text{out}}$ and $\mathbf{d}_t^{\text{out}}$ in target domain $\mathcal{T}_t$ are obtained by processing the original embeddings $\mathbf{o}_t^{\text{orig}}$ and $\mathbf{d}_t^{\text{orig}}$ through operations analogous to those described in (3) and (4), involving interactions with the transferred prompts.

Prompt-to-Representation Mapping: We propose a variant that leverages the cross-domain information characteristics of $\hat{\mathbf{p}}^o$, which are learned from the overlapping set $\mathcal{O}$. This approach is informed by knowledge transfer theory in cognitive science, suggesting that cognitive representations in one domain can be mapped to another through an intermediate transformation. Specifically, we hypothesize that a student's latent ability in one subject can be transferred to another via personalized prompts acting as intermediaries.

Personalized prompts, derived from subjects such as Physics or Biology, implicitly encode domain-general cognitive traits (e.g., logical reasoning, analytical skills). Our strategy employs a linear layer to learn the mapping relationship between these personalized prompts and shallow representations of entities in the target domains

$$\mathbf{o}_t^{\text{orig}} = \text{Linear}_{\text{init}}(\hat{\mathbf{p}}^o) \quad (7)$$

where $\mathbf{o}_t^{\text{orig}}$ is original representations from $\mathcal{O}$ in target domain $\mathcal{T}_t$. $\text{Linear}_{\text{init}}$ comprises trainable parameters that are optimized using the interaction data in the target domains. This strategy supersedes random initialization, leveraging available data to access potential original information.

By decoding these domain-general cognitive traits, we can estimate the student's capability in a target subject, enabling predictions via a shared latent cognitive space. This integration of forward mapping and reverse inference enhances the effectiveness of our model in transferring knowledge across domains.

IV. PROPOSED INSTANTIATIONS

Based on the unified framework above, we instantiate specific scenarios to illustrate its application in the following two scenarios. The pseudocodes for the proposed PromptCD-S and PromptCD-E instantiations are detailed in Algorithms 2 and 3 in the Appendix.

A. Student-Aspect CDCD: PromptCD-S

Student-aspect CDCD focuses on a cross-school scenario where $\mathcal{O}$ and $\mathcal{D}$ respectively represent a set of exercises and students. This indicates that the source and target domains have overlapping students. Each exercise item has a specific personalized prompt p_{exer}^o . We use personalized prompts to uncover the basic requirements of exercise, which are the same for students from different schools.

In this scenario, each school has a corresponding shared prompt p_{sch}^d , which reflects the collective performance of students from the common school. Then we employ the PromptCD to concatenate representations of students and exercises with their prompts in (3) and input them into the cognitive diagnosis model for interaction. By optimizing the cross-entropy loss derived from predicted scores and ground-truth from the source domain response records, PromptCD updates the prompts to extract information across source domains.

To accomplish the source-to-target prompt transfer, we concatenate p_{sch}^d for different schools in the source domains and map them to the original dimension in (6) to capture commonalities. Finally, a few records from the target domains are used to fine-tune $\hat{p}_{\text{sch}}^d$ and personalized prompts $\hat{p}_{\text{exer}}^o$ trained from source domains, enhancing their accuracy for future predictions.

B. Exercise-Aspect CDCD: PromptCD-E

Exercise-aspect CDCD, on the other hand, addresses the cross-subject scenario where $\mathcal{O}$ and $\mathcal{D}$ respectively represent a set of students and exercises. Each student is associated with a personalized prompt p_{stu}^o to capture their basic capability across various subjects, such as mathematics or physics.

Similarly, each subject has a shared prompt p_{sub}^d for all exercise items to enhance the understanding of subject-specific knowledge. Employing the pre-training and fine-tuning methodology analogous to PromptCD-S, we derive the ultimate representations of prompts.

In different scenarios, the meaning and dimensions of entity representations often differ. PromptCD mitigates the sensitivity of existing studies [21], [22] to cross-domain data by employing specific prompts to transfer information across domains, thereby aiding the cognitive diagnosis model in predicting scores accurately.

V. EXPERIMENTS

To validate the effectiveness of the PromptCD in cross-domain scenarios, we conducted extensive experiments on real-world datasets, to address the following questions.

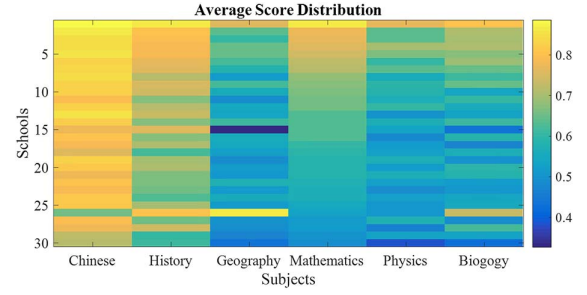


Fig. 5. Average score distribution across different schools and subjects in the SLP dataset.

- 1) *RQ1*: How does PromptCD perform in student- and exercise-aspect CDCD scenarios?
- 2) *RQ2*: How efficient are the key components in PromptCD?
- 3) *RQ3*: Can feature visualization demonstrate the effectiveness of prompts in enhancing cross-domain representations?
- 4) *RQ4*: How to conduct personalized learning guidance using PromptCD?

A. Experimental Settings

We present the experimental setup, including the datasets, baselines, metrics, and implementation details.

1) *Datasets*: SLP [20] is a real-world educational dataset that collects students from different schools' responses to multiple subjects in K-12 education. The average score distribution across different schools and subjects in the SLP dataset is illustrated in Fig. 5. Research has identified variations in the interaction levels of students when answering exercises across different subjects and schools. For students within the same school, there are similar patterns in their interaction levels across different subjects. From the perspective of student interactions, different schools represent distinct domains in the aforementioned phenomenon. This observation precisely confirms the challenges described in Section I regarding CDCD task, as similar challenges also arise from the exercise-aspect perspective.

To validate the exercise-aspect CDCD, we extracted interaction data from the SLP dataset for three humanities subjects (Chinese, History, Geography) and three science subjects (Mathematics, Physics, and Biology). For the student-aspect CDCD, we focused on interaction data from students at 30 schools in a single subject (e.g. Mathematics). We addressed the challenge of varying average cognitive levels by categorizing the schools into four bins (A, B, C, and D) based on average scores. The dataset statistics for these scenarios are presented in Table I.

2) *Baselines*: We utilized four widely recognized CD models as the backbone diagnostic models: IRT [29], MIRT [19], NeuralCD [15], and KSCD [3]. We applied our framework to these models, denoted as [Backbone]-Ours. If the prompt-to-representation mapping (Section III-C) is incorporated, the model is denoted as [Backbone]-Ours+. The original

TABLE I
DATASET STATISTICS IN THE EXPERIMENTS

Scenarios Domains	Exercise-Aspect (Humanities)			Exercise-Aspect (Sciences)			Student-Aspect (Mathematics)			
	Chinese	History	Geography	Mathematics	Physics	Biology	A-bin	B-bin	C-bin	D-bin
Student Number	4021	4021	4021	4021	4021	4021	1758	984	824	455
Exercise Number	92	164	117	137	115	120	137	137	137	137
Concept Number	14	12	24	31	34	16	31	31	31	31
Total Interactions	263 485	583 334	381 772	435 797	387 535	400 858	197 048	103 852	86 940	47 957
Interactions Per Student	66	145	95	108	96	100	112	106	106	105
Sparsity	0.29	0.12	0.19	0.21	0.16	0.17	0.18	0.23	0.23	0.23
Positive_Negative Ratio	7.04	3.03	1.69	2.80	1.49	1.96	4.46	2.72	1.87	1.35

backbone versions without cross-domain prompt transfer are used as baselines, denoted as [Backbone]-Origin. Additionally, we included state-of-the-art cross-domain cognitive diagnosis models, TechCD [21], ZeroCD [22], and CCLMF [30] as baselines, referred to as [Backbone]-Tech, [Backbone]-Zero, and [Backbone]-CCLMF, respectively. For student-aspect, TechCD and ZeroCD utilize source domain representations as initial representations for the target domains.

- 1) *Origin*: The backbones without cross-domain prompt and parameter transfer.
- 2) *Tech*: Apply TechCD [21] to the backbone. It requires the relations between knowledge concepts, constructed using a statistical method proposed in RCD [14].
- 3) *Zero*: Apply ZeroCD [22] to the backbone.
- 4) *CCLMF*: Apply CCLMF [30] to the backbone.
- 5) *Ours*: Apply the proposed PromptCD to the backbone.
- 6) *Ours+*: Using the prompt-to-representation mapping strategy on the basis of Ours.

3) *Metrics*: Cognitive diagnosis is to assess students' proficiency. However, since proficiency is an unobservable variable, researchers typically predict students' future responses (correct or incorrect) and evaluate the model's performance using the accuracy of these predictions. Therefore, we use classification evaluation metrics, namely AUC, ACC, RMSE, and F1.

4) *Implementation Details*: For both exercise- and student-aspect CDCD, we adopted a similar approach to obtain the data. In the exercise-aspect CDCD, we conducted separate diagnoses for the humanities and sciences, using any two subjects within each category as source domains and the remaining subject as the target domain. For the student-aspect CDCD, students from any three bins are treated as source domains, with the remaining bin as the target domain. In both cases, 20% of the interaction records in the target domain are randomly selected for fine-tuning, while the remaining records are used for testing. We determined the appropriate prompt dimensions for the backbone. Specifically, we set the dimensions for IRT and MIRT to 5 and 10, respectively, while NeuralCD and KSCD were set to 20.

B. Overall Comparison (RQ1)

We compare the overall performance of the models under both the student- and exercise-aspect CDCD.

Exercise-Aspect CDCD: Table II presents the performance under the exercise-aspect CDCD setting, where the bold entries indicate the best performance among different cross-domain methods under the same backbone setting (the same rule applies

to Tables III and IV). The cross-subject scenario divides six subjects into two categories: humanities and sciences, with three subjects in each category. Two subjects are treated as source domains, and one as the target domain. Across all target domains, the proposed PromptCD consistently outperforms state-of-the-art baselines. Specifically, when comparing the baselines and our framework across the IRT, MIRT, NeuralCD, and KSCD backbones, our method demonstrates significant improvement in the AUC metric, with an average increase of up to nearly 20% compared with the origin version, especially in the Biology ($0.667 \rightarrow 0.798$ in IRT) and Chinese subjects ($0.736 \rightarrow 0.864$ in IRT). Additionally, as a unified framework for cross-domain scenarios, our approach demonstrates notable improvements over TechCD [21], ZeroCD [22], and CCLMF [30]. Similarly, the RMSE metric shows a significant reduction in all scenarios. The Ours+ version, which incorporates an additional adaptation strategy, shows consistent improvements over the Ours version in most test scenarios, further validating the robustness of our approach.

Student-Aspect CDCD: Table III showcases the performance under the student-aspect CDCD. In this setup, schools are categorized into four bins (A, B, C, D) based on their average scores, with three bins designated as source domains and the remaining as the target domain. Specifically, in the A-bin target domain, the NeuralCD-Ours shows 27.9% improvement ($0.687 \rightarrow 0.879$) in the AUC metric relative to the origin, while KSCD-Ours achieves 14.5% increase ($0.764 \rightarrow 0.875$). Through the aforementioned experiments, we observed similar patterns across other metrics when applying PromptCD in various scenarios. Comparing the two tables, we see that our proposed framework shows a significant improvement over TechCD and ZeroCD in Table III than in Table II, primarily because the baseline algorithms are not specifically designed to address student-aspect CDCD.

Significance Analysis of Model Performance: We conducted Nemenyi tests on various baseline models used in different scenarios to report statistical significance for metrics such as AUC, ACC, RMSE and F1, as shown in Fig. 6. In the figure, models connected by the same horizontal line do not show statistically significant differences, while models located on different lines exhibit significant performance differences at the significance level of $\alpha = 0.05$. The PromptCD model (especially the "Ours+" version) significantly outperforms other baseline models across multiple domains. This further demonstrates the effectiveness and robustness of cross-domain prompt transfer methods in cognitive diagnosis tasks.

TABLE II
COMPARISON RESULTS IN EXERCISE-ASPECT CDCD SCENARIOS

Target Metrics	Biology				Mathematics				Physics			
	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1
IRT-Origin	0.667	0.652	0.466	0.738	0.699	0.725	0.434	0.814	0.761	0.706	0.442	0.757
IRT-Tech	0.779	0.729	0.421	0.800	0.861	0.819	0.357	0.885	0.845	0.767	0.397	0.808
IRT-Zero	0.721	0.696	0.460	0.809	0.809	0.788	0.414	0.871	0.805	0.736	0.421	0.799
IRT-CCLMF	0.775	0.723	0.425	0.808	0.870	0.817	0.353	0.884	0.850	0.761	0.403	0.813
IRT-Ours	0.798	0.742	0.411	0.815	0.874	0.826	0.347	0.886	0.858	0.776	0.387	0.818
IRT-Ours+	0.799	0.743	0.411	0.816	0.880	0.832	0.342	0.890	0.864	0.781	0.384	0.823
MIRT-Origin	0.672	0.669	0.459	0.763	0.712	0.746	0.421	0.835	0.771	0.717	0.440	0.774
MIRT-Tech	0.779	0.727	0.422	0.797	0.864	0.817	0.355	0.878	0.847	0.766	0.395	0.807
MIRT-Zero	0.723	0.706	0.439	0.799	0.786	0.786	0.395	0.867	0.802	0.734	0.420	0.781
MIRT-CCLMF	0.771	0.726	0.426	0.805	0.867	0.810	0.368	0.879	0.843	0.768	0.398	0.813
MIRT-Ours	0.793	0.738	0.415	0.816	0.881	0.834	0.347	0.893	0.855	0.775	0.398	0.822
MIRT-Ours+	0.801	0.743	0.411	0.809	0.886	0.838	0.343	0.893	0.865	0.785	0.389	0.821
NCDM-Origin	0.706	0.655	0.456	0.724	0.755	0.775	0.403	0.856	0.790	0.725	0.432	0.771
NCDM-Tech	0.780	0.727	0.421	0.795	0.865	0.809	0.360	0.874	0.797	0.732	0.423	0.784
NCDM-Zero	0.734	0.697	0.437	0.792	0.824	0.788	0.373	0.871	0.791	0.714	0.438	0.746
NCDM-CCLMF	0.765	0.731	0.424	0.811	0.844	0.808	0.363	0.875	0.839	0.769	0.403	0.809
NCDM-Ours	0.785	0.735	0.417	0.815	0.852	0.813	0.359	0.878	0.848	0.764	0.397	0.796
NCDM-Ours+	0.788	0.731	0.418	0.812	0.872	0.816	0.357	0.874	0.861	0.782	0.386	0.820
KSCD-Origin	0.710	0.691	0.445	0.779	0.761	0.774	0.401	0.854	0.797	0.729	0.426	0.772
KSCD-Tech	0.778	0.729	0.422	0.799	0.859	0.818	0.356	0.882	0.842	0.765	0.398	0.807
KSCD-Zero	0.728	0.703	0.431	0.792	0.801	0.793	0.382	0.872	0.798	0.732	0.428	0.788
KSCD-CCLMF	0.782	0.732	0.420	0.799	0.861	0.815	0.357	0.879	0.843	0.765	0.397	0.813
KSCD-Ours	0.795	0.741	0.413	0.818	0.869	0.826	0.349	0.888	0.855	0.777	0.389	0.817
KSCD-Ours+	0.796	0.739	0.414	0.812	0.870	0.828	0.350	0.886	0.855	0.776	0.389	0.816
Target Metrics	Chinese				History				Geography			
	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1
IRT-Origin	0.736	0.872	0.319	0.931	0.707	0.752	0.415	0.843	0.687	0.659	0.465	0.731
IRT-Tech	0.830	0.862	0.319	0.922	0.799	0.735	0.417	0.810	0.754	0.707	0.438	0.782
IRT-Zero	0.821	0.878	0.306	0.934	0.790	0.774	0.391	0.859	0.760	0.705	0.437	0.780
IRT-CCLMF	0.855	0.872	0.296	0.934	0.803	0.778	0.381	0.865	0.776	0.711	0.433	0.786
IRT-Ours	0.864	0.886	0.290	0.937	0.811	0.791	0.377	0.870	0.783	0.722	0.425	0.790
IRT-Ours+	0.865	0.886	0.289	0.938	0.814	0.791	0.376	0.871	0.787	0.724	0.424	0.793
MIRT-Origin	0.744	0.873	0.319	0.932	0.719	0.756	0.410	0.846	0.689	0.671	0.461	0.763
MIRT-Tech	0.811	0.742	0.411	0.835	0.795	0.752	0.408	0.829	0.761	0.709	0.435	0.781
MIRT-Zero	0.822	0.879	0.302	0.935	0.782	0.776	0.392	0.856	0.724	0.662	0.454	0.724
MIRT-CCLMF	0.845	0.867	0.296	0.930	0.804	0.774	0.385	0.861	0.758	0.704	0.439	0.776
MIRT-Ours	0.863	0.885	0.289	0.937	0.818	0.791	0.376	0.871	0.778	0.714	0.431	0.793
MIRT-Ours+	0.866	0.888	0.288	0.937	0.822	0.794	0.374	0.870	0.792	0.728	0.422	0.787
NCDM-Origin	0.782	0.851	0.323	0.916	0.742	0.740	0.412	0.826	0.717	0.679	0.453	0.752
NCDM-Tech	0.809	0.875	0.308	0.933	0.740	0.769	0.405	0.866	0.721	0.678	0.451	0.749
NCDM-Zero	0.805	0.876	0.306	0.932	0.742	0.761	0.407	0.863	0.728	0.687	0.445	0.755
NCDM-CCLMF	0.838	0.865	0.305	0.930	0.789	0.772	0.397	0.850	0.766	0.711	0.438	0.775
NCDM-Ours	0.843	0.881	0.297	0.934	0.793	0.778	0.389	0.856	0.774	0.717	0.430	0.780
NCDM-Ours+	0.852	0.878	0.296	0.931	0.807	0.788	0.379	0.868	0.786	0.720	0.426	0.783
KSCD-Origin	0.787	0.869	0.318	0.929	0.744	0.771	0.402	0.857	0.722	0.688	0.448	0.771
KSCD-Tech	0.829	0.797	0.371	0.875	0.798	0.749	0.410	0.825	0.756	0.706	0.438	0.772
KSCD-Zero	0.825	0.851	0.328	0.914	0.789	0.769	0.406	0.853	0.749	0.702	0.443	0.778
KSCD-CCLMF	0.842	0.875	0.305	0.925	0.793	0.776	0.388	0.856	0.772	0.711	0.432	0.787
KSCD-Ours	0.854	0.883	0.292	0.936	0.807	0.789	0.380	0.868	0.785	0.723	0.425	0.794
KSCD-Ours+	0.854	0.884	0.292	0.936	0.804	0.789	0.380	0.869	0.787	0.719	0.426	0.797

Note: The bold entries indicate the best performance among different cross-domain methods under the same backbone setting.

C. Detailed Analysis (RQ2)

To address the role of PromptCD in key aspects, we analyze the effects of different modules in the proposed model.

Different Fine-Tuning Ratios: We examine the performance of the PromptCD across varying proportions of fine-tuning data, addressing the challenge of data sparsity in cross-domain knowledge adaptation. Specifically, we evaluate the model's performance with fine-tuning data proportions of 0.1, 0.2, and 0.3 on both the exercise-aspect CDCD for the biology target domain and the student-aspect for the A-bin target domain, shown in Fig. 7(a) and 7(b). In most cases, the performance improves as the proportion increases, even at a small proportion of 0.1, the models can achieve satisfactory performance in dual-aspect scenarios by using PromptCD.

Different Prompt Dimensions: We demonstrate the influence of prompt dimensionality on PromptCD. We provide results

using the NeuralCD, which employs horizontal concatenation, as an example. Specifically, we set the prompt dimensionality to 1, 5, 10, or 20. The comparison results are shown in Fig. 7(c) and 7(d). Performance is lowest when the prompt dimension is 1. As the prompt dimension increases, performance improves, reflecting the enhanced information capacity of the prompts and their effectiveness in improving model performance.

Various Source Domains: We evaluate the impact of various source domains PromptCD. Specifically, for the exercise-aspect scenario with biology as the target domain, we consider three scenarios: using mathematics, physics, or both as source domains, respectively. Similarly, for the student-aspect scenario with A-bin as the target domain, we consider three types of combinations of source domains, which are illustrated in Table IV. The performance achieved using data from two source domains is superior to that obtained from a single source

TABLE III
COMPARISON RESULTS IN STUDENT-ASPECT CDCD SCENARIOS

Target Metrics	A-bin				B-bin				C-bin				D-bin			
	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1
IRT-Origin	0.670	0.800	0.393	0.884	0.548	0.728	0.451	0.841	0.678	0.679	0.477	0.777	0.519	0.537	0.512	0.630
IRT-Tech	0.821	0.847	0.336	0.912	0.837	0.811	0.365	0.877	0.769	0.741	0.425	0.830	0.809	0.733	0.419	0.767
IRT-Zero	0.855	0.854	0.324	0.917	0.854	0.823	0.360	0.884	0.847	0.784	0.390	0.835	0.831	0.753	0.406	0.792
IRT-CCLMF	0.851	0.854	0.321	0.920	0.853	0.822	0.362	0.883	0.855	0.790	0.384	0.847	0.839	0.757	0.403	0.799
IRT-Ours	0.871	0.867	0.316	0.922	0.870	0.826	0.351	0.884	0.877	0.806	0.368	0.858	0.857	0.766	0.397	0.814
IRT-Ours+	0.881	0.872	0.308	0.925	0.881	0.834	0.344	0.89	0.881	0.811	0.363	0.858	0.858	0.764	0.395	0.813
MIRT-Origin	0.718	0.820	0.372	0.896	0.742	0.762	0.413	0.848	0.715	0.706	0.461	0.798	0.729	0.682	0.470	0.749
MIRT-Tech	0.820	0.845	0.339	0.912	0.840	0.811	0.361	0.878	0.805	0.760	0.403	0.831	0.817	0.739	0.414	0.780
MIRT-Zero	0.841	0.841	0.347	0.902	0.849	0.814	0.366	0.880	0.845	0.799	0.367	0.854	0.804	0.735	0.433	0.754
MIRT-CCLMF	0.834	0.849	0.331	0.911	0.842	0.814	0.368	0.881	0.842	0.770	0.398	0.841	0.814	0.742	0.428	0.762
MIRT-Ours	0.861	0.859	0.324	0.917	0.863	0.816	0.365	0.883	0.862	0.778	0.395	0.844	0.844	0.766	0.409	0.806
MIRT-Ours+	0.886	0.872	0.311	0.923	0.886	0.836	0.347	0.891	0.881	0.807	0.375	0.858	0.859	0.778	0.398	0.813
NCDM-Origin	0.687	0.809	0.387	0.894	0.693	0.743	0.465	0.847	0.709	0.643	0.485	0.782	0.533	0.574	0.653	0.729
NCDM-Tech	0.818	0.845	0.340	0.910	0.838	0.805	0.364	0.873	0.816	0.766	0.401	0.835	0.796	0.697	0.429	0.714
NCDM-Zero	0.761	0.800	0.376	0.877	0.798	0.712	0.430	0.783	0.754	0.724	0.430	0.804	0.797	0.706	0.455	0.718
NCDM-CCLMF	0.844	0.851	0.332	0.911	0.846	0.813	0.368	0.871	0.838	0.775	0.396	0.820	0.813	0.734	0.414	0.757
NCDM-Ours	0.879	0.870	0.308	0.924	0.865	0.820	0.357	0.877	0.856	0.780	0.390	0.822	0.837	0.751	0.408	0.765
NCDM-Ours+	0.878	0.865	0.319	0.920	0.878	0.833	0.352	0.891	0.864	0.796	0.380	0.841	0.840	0.752	0.409	0.785
KSCD-Origin	0.764	0.831	0.368	0.901	0.756	0.775	0.414	0.860	0.766	0.726	0.448	0.807	0.769	0.706	0.449	0.737
KSCD-Tech	0.809	0.848	0.338	0.913	0.839	0.809	0.366	0.884	0.780	0.756	0.412	0.830	0.794	0.725	0.425	0.772
KSCD-Zero	0.778	0.806	0.346	0.877	0.808	0.798	0.390	0.879	0.785	0.763	0.408	0.839	0.797	0.737	0.423	0.782
KSCD-CCLMF	0.854	0.855	0.327	0.915	0.853	0.822	0.358	0.886	0.853	0.794	0.385	0.846	0.818	0.759	0.418	0.786
KSCD-Ours	0.875	0.865	0.312	0.921	0.876	0.834	0.346	0.891	0.871	0.806	0.371	0.854	0.844	0.768	0.403	0.805
KSCD-Ours+	0.880	0.868	0.310	0.923	0.878	0.835	0.345	0.891	0.870	0.804	0.371	0.851	0.846	0.771	0.404	0.795

Note: The bold entries indicate the best performance among different cross-domain methods under the same backbone setting.

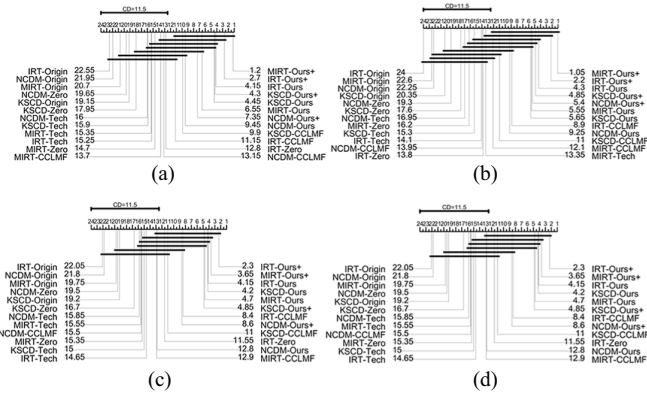


Fig. 6. Visualization of significance test results for evaluation metrics. Models connected by the same horizontal line do not show statistically significant differences, while models located on different lines exhibit significant performance differences at the significance level of $\alpha = 0.05$. (a) ACC significance test. (b) AUC significance test. (c) RMSE significance test. (d) F1 significance test.

domain, which indicates that PromptCD can learn common knowledge across multiple source domains.

Various Cross-Domain Types: Table V presents the experimental results of the model in more scenarios, including from humanities to sciences and from sciences to humanities. The detailed experimental setups are as follows: Sciences-Humanities: Source: Biology, Mathematics $\rightarrow$ Target: Geography. Humanities-Humanities: Source: Chinese, History $\rightarrow$ Target: Geography. Humanities-Sciences: Source: Chinese, History $\rightarrow$ Target: Physics. Sciences-Sciences: Source: Biology, Mathematics $\rightarrow$ Target: Physics. Our results indicate that in both cases, PromptCD outperforms the comparison algorithms. Interestingly, its performance shows a slight decline compared with the “humanities to sciences” and “sciences to

sciences” scenarios, which aligns with empirical expectations. Specifically, the transfer of student states is more effective between disciplines with similar characteristics, whereas greater differences between disciplines result in more significant deviations in the transferred student states.

Alternative Architectures for Prompt Fusion: To further validate our architectural choices, we conducted additional ablation experiments to explore alternative prompt fusion strategies beyond simple concatenation followed by a linear transformation. Specifically, we evaluate the following variants: **MLP Fusion:** Applying a multilayer perceptron (MLP) instead of Linear_{s2t} to capture nonlinear interactions between prompts. **Element-Wise Addition:** The original embedding and the corresponding prompt features are added element-wise to generate the final individual representation. The new features effectively integrate the local attributes of the node with the global information carried by the prompt. **Element-Wise Multiplication:** The original embedding and the corresponding prompt features are multiplied element-wise to generate the final individual representation. By applying element-wise multiplication, the global semantic information of the prompt features can be enhanced while preserving important detailed information. These fusion strategies are evaluated under both student-aspect and exercise-aspect CDCD settings, with performance compared in terms of AUC, ACC, and generalization capability. The results in Table VI demonstrate that our method achieves superior performance across multiple metrics.

Different Types of Prompts: We analyze the impact of removing different prompt components by evaluating three variants: **MIRT-1:** Removes personalized prompts $\hat{p}^o$ to assess their contribution. **MIRT-2:** Removes shared prompts $\hat{p}_t^d$ to evaluate their role in cross-domain adaptation. **MIRT-3:** Removes both components to observe the overall effect. The results, shown in Table VII, confirm that both personalized and shared prompts

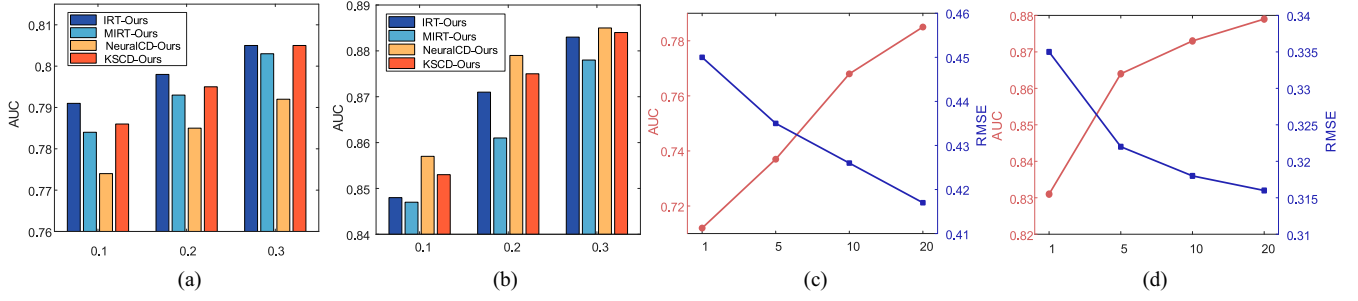


Fig. 7. Performance comparisons with (a) and (b) different tuning ratios and (c) and (d) different prompt dimensions.

TABLE IV
PERFORMANCE COMPARISONS WITH DIFFERENT SOURCE DOMAINS IN EXERCISE AND STUDENT-ASPECT SCENARIOS

Backbone	Source Domain	Biology				Source Domain	A-bin			
		AUC	ACC	RMSE	F1		AUC	ACC	RMSE	F1
IRT	Mathematics	0.781	0.729	0.414	0.814	B	0.860	0.861	0.321	0.919
	Physics	0.782	0.729	0.413	0.814	B+C	0.870	0.865	0.314	0.921
	Mathematics+Physics	0.798	0.742	0.411	0.815	B+C+D	0.871	0.867	0.316	0.922
MIRT	Mathematics	0.781	0.730	0.421	0.814	B	0.853	0.856	0.324	0.917
	Physics	0.784	0.732	0.419	0.814	B+C	0.862	0.861	0.320	0.919
	Mathematics+Physics	0.793	0.738	0.415	0.816	B+C+D	0.861	0.859	0.321	0.919
NeuralCD	Mathematics	0.774	0.726	0.423	0.796	B	0.865	0.862	0.318	0.920
	Physics	0.773	0.726	0.423	0.800	B+C	0.874	0.868	0.311	0.922
	Mathematics+Physics	0.785	0.735	0.417	0.815	B+C+D	0.879	0.870	0.308	0.924
KSCD	Mathematics	0.782	0.735	0.418	0.812	B	0.868	0.864	0.316	0.921
	Physics	0.781	0.734	0.418	0.815	B+C	0.875	0.868	0.311	0.923
	Mathematics+Physics	0.795	0.741	0.413	0.818	B+C+D	0.875	0.865	0.312	0.921

Note: The bold entries indicate the best performance among different cross-domain methods under the same backbone setting.

TABLE V
COMPARISON RESULTS OF VARIOUS CDCD SCENARIOS

Scenarios Metrics	Sciences-Humanities				Humanities-Humanities				Humanities-Sciences				Sciences-Sciences			
	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1	AUC	ACC	RMSE	F1
IRT-Origin	0.677	0.650	0.472	0.721	0.687	0.659	0.465	0.731	0.747	0.686	0.460	0.731	0.761	0.706	0.442	0.757
IRT-Tech	0.757	0.705	0.438	0.784	0.754	0.707	0.438	0.782	0.829	0.753	0.407	0.791	0.845	0.767	0.397	0.808
IRT-Zero	0.734	0.693	0.446	0.776	0.760	0.705	0.437	0.780	0.804	0.730	0.424	0.766	0.805	0.736	0.421	0.799
IRT-CCLMF	0.775	0.715	0.427	0.791	0.776	0.711	0.433	0.786	0.804	0.730	0.424	0.766	0.805	0.736	0.421	0.799
IRT-Ours	0.791	0.726	0.422	0.795	0.783	0.722	0.425	0.790	0.854	0.773	0.390	0.815	0.858	0.776	0.387	0.818
IRT-Ours+	0.792	0.727	0.421	0.792	0.787	0.724	0.424	0.793	0.854	0.774	0.390	0.814	0.864	0.781	0.384	0.823
MIRT-Origin	0.695	0.669	0.459	0.773	0.689	0.671	0.461	0.763	0.775	0.718	0.439	0.779	0.771	0.717	0.440	0.774
MIRT-Tech	0.765	0.709	0.435	0.787	0.761	0.709	0.435	0.781	0.828	0.744	0.416	0.807	0.847	0.766	0.395	0.807
MIRT-Zero	0.722	0.686	0.450	0.763	0.724	0.662	0.454	0.724	0.794	0.721	0.428	0.757	0.802	0.734	0.420	0.781
MIRT-CCLMF	0.767	0.710	0.434	0.775	0.758	0.704	0.439	0.776	0.838	0.756	0.399	0.810	0.843	0.768	0.398	0.813
MIRT-Ours	0.786	0.720	0.428	0.794	0.778	0.714	0.431	0.793	0.843	0.764	0.405	0.813	0.855	0.775	0.398	0.822
MIRT-Ours+	0.791	0.728	0.423	0.789	0.792	0.728	0.422	0.787	0.852	0.771	0.397	0.806	0.865	0.785	0.389	0.821
NCDM-Origin	0.714	0.683	0.460	0.765	0.717	0.679	0.453	0.752	0.782	0.721	0.435	0.764	0.790	0.725	0.432	0.771
NCDM-Tech	0.771	0.715	0.431	0.788	0.721	0.678	0.451	0.749	0.821	0.743	0.414	0.799	0.797	0.732	0.423	0.784
NCDM-Zero	0.718	0.688	0.451	0.777	0.728	0.687	0.445	0.755	0.798	0.729	0.426	0.787	0.791	0.714	0.438	0.746
NCDM-CCLMF	0.769	0.715	0.435	0.781	0.766	0.711	0.438	0.775	0.831	0.754	0.408	0.791	0.839	0.769	0.403	0.809
NCDM-Ours	0.779	0.720	0.427	0.784	0.774	0.717	0.430	0.780	0.832	0.757	0.407	0.806	0.848	0.764	0.397	0.796
NCDM-Ours+	0.786	0.722	0.424	0.783	0.786	0.720	0.426	0.783	0.848	0.769	0.397	0.814	0.861	0.782	0.386	0.820
KSCD-Origin	0.722	0.687	0.448	0.770	0.722	0.688	0.448	0.771	0.798	0.728	0.426	0.771	0.797	0.729	0.426	0.772
KSCD-Tech	0.761	0.709	0.435	0.785	0.756	0.706	0.438	0.772	0.826	0.753	0.409	0.805	0.842	0.765	0.398	0.807
KSCD-Zero	0.734	0.695	0.442	0.776	0.749	0.702	0.443	0.778	0.809	0.733	0.420	0.791	0.798	0.732	0.428	0.788
KSCD-CCLMF	0.776	0.719	0.428	0.796	0.772	0.711	0.432	0.787	0.836	0.761	0.404	0.806	0.843	0.765	0.397	0.813
KSCD-Ours	0.788	0.726	0.424	0.791	0.785	0.723	0.425	0.794	0.848	0.769	0.395	0.804	0.855	0.777	0.389	0.817
KSCD-Ours+	0.788	0.724	0.424	0.794	0.787	0.719	0.426	0.797	0.849	0.772	0.393	0.815	0.855	0.776	0.389	0.816

Note: The bold entries indicate the best performance among different cross-domain methods under the same backbone setting.

contribute significantly to performance. Removing either component leads to a noticeable decline, highlighting their complementary roles in capturing individualized and domain-level knowledge.

D. Feature Visualization (RQ3)

In this section, we visualized the representations learned by the model in CDCD scenarios for both exercises and students

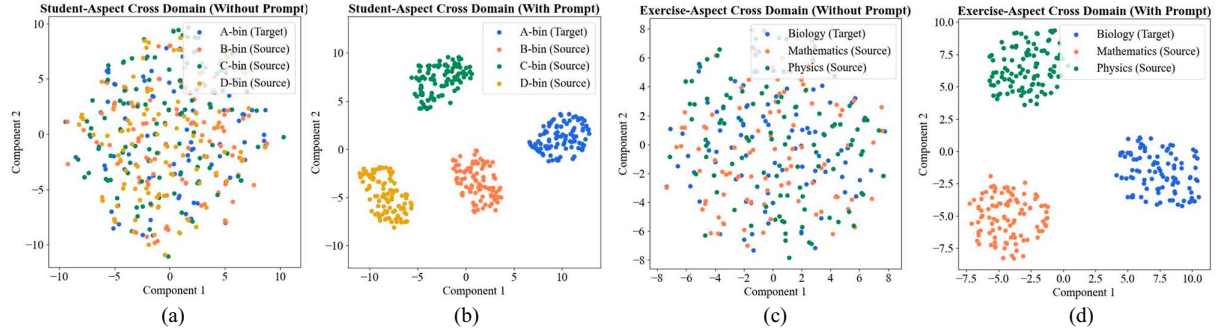


Fig. 8. Visualization of (a) and (c) origin representations without prompt and (b) and (d) our representations with prompt.

TABLE VI
PERFORMANCE COMPARISON OF DIFFERENT PROMPT
FUSION STRATEGIES

Target	Mathematics			
	AUC	ACC	RMSE	F1
MIRT-Ours+	0.886	0.838	0.343	0.893
MIRT-MLP	0.881	0.833	0.350	0.890
MIRT-Addition	0.882	0.833	0.347	0.890
MIRT-Multiplication	0.744	0.754	0.410	0.850

Target	A-bin			
	AUC	ACC	RMSE	F1
MIRT-Ours+	0.886	0.872	0.311	0.923
MIRT-MLP	0.877	0.865	0.318	0.920
MIRT-Addition	0.881	0.853	0.322	0.916
MIRT-Multiplication	0.769	0.816	0.389	0.898

Note: The bold entries indicate the best performance among different cross-domain methods under the same backbone setting.

TABLE VII
ABLATION STUDY ON THE IMPACT OF PROMPT
COMPONENTS

Target	Mathematics			
	AUC	ACC	RMSE	F1
MIRT-Ours+	0.886	0.838	0.343	0.893
MIRT-1	0.728	0.758	0.410	0.852
MIRT-2	0.879	0.830	0.351	0.890
MIRT-3	0.712	0.746	0.421	0.835

Target	A-bin			
	AUC	ACC	RMSE	F1
MIRT-Ours+	0.886	0.872	0.311	0.923
MIRT-1	0.732	0.831	0.359	0.904
MIRT-2	0.877	0.865	0.317	0.920
MIRT-3	0.718	0.820	0.372	0.896

Note: The bold entries indicate the best performance among different cross-domain methods under the same backbone setting.

to explain why the prompts learned from the source domains are effective.

For the student-aspect CDCD scenario, Fig. 8(a) displays the distribution of the original student representations after dimensionality reduction and reveals that the original representations of students from different bins do not display any discernible patterns. In contrast, Fig. 8(b) displays the distribution of the final representations, obtained after the operations described in Section III-B. The clustering of students within the same bin suggests that our prompts effectively capture the overall characteristics of the domain. Additionally, student representations from A-bin are more distant from those of D-bin and closer to those of B-bin, indicating that the prompts effectively distinguish between different levels of student groups. In the exercise-aspect CDCD, as shown in Fig. 8(c) and 8(d), the

TABLE VIII
CLUSTER ANALYSIS BEFORE AND AFTER PROMPTING

Status	Dimension	Intra-Dist	Inter-Dist
without Prompt	Exercise Embedding	6.6352	0.2911
without Prompt	Student Embedding	7.8856	0.6041
with Prompt	Exercise Embedding	2.8690	12.2945
with Prompt	Student Embedding	2.9120	13.0087

transfer prompts also effectively capture both the internal characteristics of each subject and the distinctions between different subjects.

Additionally, we employed two quantitative metrics—intercluster distance and intracluster distance—to analyze the changes in student and exercise representations before and after introducing prompts. As shown in Table VIII, the intercluster and intracluster distances are significantly smaller after introducing prompts compared with before, providing a clearer demonstration of the effectiveness of our method and ensuring consistency in the presentation across both student and exercise dimensions.

E. Personalized Recommendation (RQ4)

In this section, we illustrate how the PromptCD facilitates personalized learning guidance for exercise recommendations. We employ a straightforward yet effective strategy to suggest exercises related to concepts that students have not yet mastered [21], ensuring they are of appropriate difficulty [5]. Based on a CD backbone, we first determine whether this is a student-side or exercise-side cross-domain recommendation scenario, then select the corresponding framework to be added to the CD model. After the model undergoes a two-stage training process, it produces a well-trained model $\mathcal{M}$. $\mathcal{M}$ can determine the student's mastery level of knowledge concepts through a diagnostic module. We select N exercises associated with the knowledge concepts that the student has not yet mastered. Furthermore, we aim for the exercises to be of moderate difficulty for the student, as exercises that are too difficult or too easy may hurt their learning interest. Therefore, from this set of K exercises, we ultimately choose K exercises of moderate difficulty to form the recommended list for the student. Let's take the example of applying the proposed framework to NeuralCD [2] to showcase the results of the personalized learning recommendation. In the exercise-aspect CDCD, mathematics and biology are treated as the source

TABLE IX
RECOMMENDING EXAMPLE IN EXERCISE-ASPECT CDCD

Concept ID	11	14	19	25	26	28	33
Exercise ID	16	2	14	25	23	54	3
Student Mastery	0.491	0.472	0.484	0.489	0.490	0.481	0.485
Exercise Difficulty	0.550	0.481	0.520	0.428	0.530	0.487	0.469
True Performance	0	0	0	1	0	0	1

domains, and physics is the target domain. We recommend exercises to a randomly sampled student. The recommended concept ID, exercise ID, corresponding student mastery, exercise difficulty, and true performance are detailed in Table IX. The results demonstrate that the recommended exercises align with the requirements of practical applications, which are exercises that the student has not mastered yet and are of moderate difficulty.

VI. RELATED WORK

A. Cognitive Diagnosis

Cognitive diagnosis [1], [3], [14] is a form of student learning modeling that plays a vital role in educational recommendation tasks [4], [5]. Traditional cognitive diagnosis models, such as IRT [29] and MIRT [19], utilize unidimensional and multidimensional latent traits, respectively, to represent student and exercise features. The DINA model [31] incorporates guessing and slipping parameters but often relies on assumptions that oversimplify student interactions. These traditional models have established a foundation for cognitive diagnosis but struggle to capture the complexities of student behavior and learning patterns.

To address these limitations, various deep learning-based cognitive diagnosis models have been proposed. Wang et al. [15] introduced NeuralCD, which uses neural networks to enhance both accuracy and interpretability. Many works have expanded upon NeuralCD, such as KaNCD [15] and KSCD [3], which make full use of information from noninteractive knowledge concepts. Beyond these, recent studies have explored more advanced cognitive diagnosis frameworks. For example, ORCDF [32] introduces an oversmoothing-resistant model to enhance learning representations in online education systems, while SCD [33] combines symbolic reasoning with hybrid optimization to improve diagnosis performance. Moreover, FineCD [34] explores how foundation models can enhance derivative-free cognitive diagnosis, improving generalization across diverse student populations and subjects. These models that assume the training and testing data are from the same distribution will experience a significant decline in performance when confronted with nonidentical distributions.

1) *Cross-Domain Cognitive Diagnosis*: The introduction of new domains in online education often leads to the unavailability of practice logs for many students, creating the CDCD issue [21]. Gao et al. [21] proposed TechCD, which uses transferable knowledge concept graphs to address the cold-start problem in new domains. This method embeds knowledge concepts and student behaviors into a graph, leveraging transferable knowledge to accurately assess cognitive abilities. ZeroCD [22] tackles the CDCD problem by utilizing early-participating student

data to assess cognitive abilities with minimal data. In addition, Hu et al. [30] proposed CCLMF, which leverages a meta-learner to predict network parameters and enhances model performance on target courses using knowledge from source courses. However, most of them only focuses on just one aspect of the issue. LRCD [35] introduces a language representation-favored approach demonstrating the potential of pretrained language models in zero-shot CDCD settings. In this paper, we focus on a scenario-agnostic CDCD framework, maintaining compatibility with both exercise-aspect and student-aspect scenarios.

B. Prompt Learning

Prompt learning [36] is a technique applied to pre-trained language models [23] and has demonstrated significant success in various applications, including recommendation tasks [37], [38], [39], [40]. This approach guides the model's generation process using prompts. Prompts can be either hard (discrete words) or soft (continuous learnable embeddings) [41]. Soft prompts, in particular, offer greater flexibility as they can be optimized and adjusted during training, allowing them to better adapt to specific tasks and data requirements [42].

1) *Prompt Learning for Cross-Domain Tasks*: The prompt learning method, through the adjustment and optimization of prompts, can adapt to the language and characteristics of different domains. Consequently, it has shown promising results in cross-domain recommendation tasks [24], [25], [26], [27], [28]. In these tasks, shared knowledge from source domains is often transferred to the target domain via knowledge-enhanced prompts. Unlike hard prompts, which require extensive hand-crafting and are highly specific to individual tasks, soft prompts offer greater flexibility and can be more easily optimized and adapted [27]. This adaptability reduces inefficiencies and improves robustness across various tasks and models. Hard prompts, on the other hand, often present significant challenges; poorly designed prompts can negatively impact model performance and may not transfer effectively across tasks [43]. Moreover, the need for prompt engineering and the difficulty of creating effective templates for each task further limit their efficiency [44], [45]. Therefore, this article focuses on leveraging soft prompt learning for the CDCD task.

VII. CONCLUSION

In this article, we proposed the PromptCD framework for cross-domain cognitive diagnosis tasks in intelligence education. Specifically, we designed the prompts to enhance the student and exercise representation across domains. The prompts follow a two-stage mode of pretraining and fine-tuning. Importantly, the proposed framework can be applied to both student-aspect and exercise-aspect cross-domain scenarios. Experimental results on the real-world datasets illustrated the effectiveness of the proposed framework.

APPENDIX

In this section, the experimental details, the pseudocodes of the proposed models, and more experiments and analysis are presented.

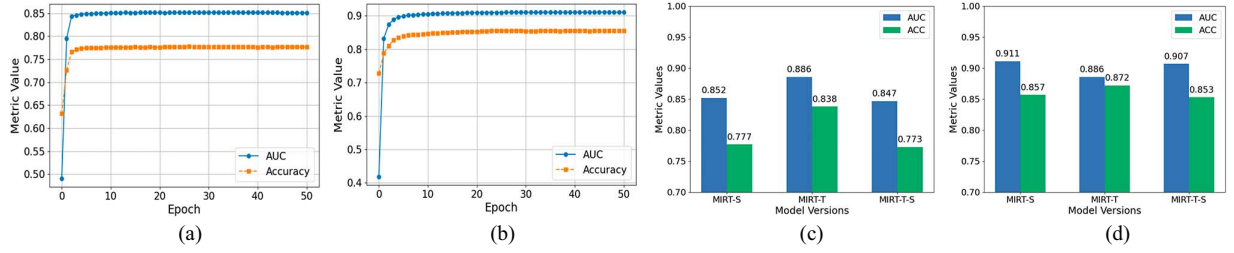


Fig. 9. Performance trends (a) and (b) and knowledge retention across phases (c) and (d).

A. Experimental Details of Section I

The experiment was conducted on the SLP dataset using the MIRT model, focusing on three subjects: Mathematics, Physics, and Chinese. The Mathematics dataset was randomly divided into three subsets—Mathematics1, Mathematics2, and Mathematics3—with a ratio of 2:2:6. In this configuration, Mathematics2 served as the testing set, while the other two subsets were combined in various ways for training. This approach allowed for diverse training sets, enhancing the robustness of the results. The hyperparameter settings for the MIRT model included a latent trait dimension of 10, a learning rate of 0.001, and a batch size of 256.

B. Pseudocodes of PromptCD-S and PromptCD-E

Pseudocodes of PromptCD-S and PromptCD-E applied to NeuralCD are shown in Algorithms 2 and 3.

C. Hyperparameter Tuning Strategies and Training Epoch Ratios in Experiments

We first set the initial hyperparameters based on prior domain knowledge and experience, typically setting the learning rate to 0.001 and the batch size to 256. Then, we conduct a learning rate range test, gradually increasing the learning rate (e.g., doubling it each time). After each adjustment, we train for a fixed number of batches and record the corresponding loss values. By plotting the relationship between the learning rate and loss, we identify the range where the loss decreases most rapidly and remains stable. Additionally, for batch size adjustments, we follow a linear scaling rule: when the batch size increases by a factor of k , the learning rate is also increased by a factor of k accordingly to ensure training stability and efficiency.

In practical experiments, we observed that datasets composed of different disciplines require varying numbers of iterations to reach convergence during the pre-training and fine-tuning stages. Therefore, using a fixed training epoch ratio may not be suitable for different scenarios. To address this issue, we introduce an Early Stopping mechanism. Specifically, during training, we continuously monitor performance metrics on the validation set (e.g., loss value or accuracy). If the performance does not improve for a certain number of consecutive iterations, training is terminated early. This strategy not only significantly reduces unnecessary computational costs but also effectively prevents overfitting, ensuring both training efficiency and model generalization capability.

Algorithm 2: PromptCD-S for NeuralCD.

```

1: Input: The cognitive diagnosis model  $\mathcal{M}$  based on the NeuralCD backbone.  $\mathcal{M}$ ,
   records  $LS$  for pre-training and  $LT_t^{few}$  in target domain  $t$  for fine-tuning.
2: Output: fine-tuned model  $\mathcal{M}$ , the transfer prompts  $\hat{p}_{exer}^o$  and  $\hat{p}_{sch}^d$ .
3: —Pre-training Stage—
4: while  $e_1 \leq Epoch_{Pretrain}$  do
5:   for  $LS_s \in \{LS_1, LS_2, \dots, LS_{|S|}\}$  do
6:     Initialize the students embedding  $\alpha_s^{orig}$ , the exercises embedding  $\beta_s^{orig}$ ,
       prompts  $p_{exer}^o, p_{sch}^d$  and  $\mathcal{M}$ ;
7:     Enhance the representation of students and exercises in Eq.(3) and Eq.(4),
       connect  $p_{exer}^o$  to  $\beta_s^{orig}$ , and  $p_{sch}^d$  to  $\alpha_s^{orig}$ , obtaining  $\alpha_s^{out}$  and  $\beta_s^{out}$  after
       mapping;
8:     Input  $\alpha_s^{out}$  and  $\beta_s^{out}$  to  $\mathcal{M}$ . Specifically, subtract  $\beta_s^{out}$  from  $\alpha_s^{out}$ , multiply by
       the exercise and knowledge vectors, and pass through NeuralCD to get the
       predicted score  $y_s$ ;
9:     Calculate the loss using  $LS_s$  to update the model;
10:   end for
11: end while
12: —Fine-Tuning Stage—
13: while  $e_2 \leq Epoch_{FineTune}$  do
14:   Initialize the students embedding  $\alpha_t^{orig}$ , the exercises embedding  $\beta_t^{orig}$ ,
15:   Obtain the transfer prompts  $\hat{p}_{exer}^o$  and  $\hat{p}_{sch}^d$  in Eq.(5) and Eq.(6);
16:   Activate improvement policy in Eq.(7);
17:   Enhance the representations in a manner similar to pre-training;
18:   Input  $\alpha_t^{out}$  and  $\beta_t^{out}$  to  $\mathcal{M}$  to obtain the final predicted score  $y_t$ , as in the
       pre-training process;
19:   Calculate the loss using  $LT_t^{few}$  to update the model.
20: end while

```

D. Convergence Analysis of Pretraining and Fine-tuning

To analyze convergence, we take the target domain as the mathematics subject in a cross-disciplinary scenario and track the AUC and ACC curves over the two training stages. As shown in Fig. 9, the experimental results demonstrate that during the pretraining phase, the loss values stabilized after sufficient iterations, providing robust initialization for subsequent fine-tuning. The fine-tuning phase further optimized model performance, and after adequate iterations, the evaluation metrics showed no significant fluctuations, indicating that the model achieved stable adaptation to the target domain.

E. Trade-Off Between Knowledge Preservation and Adaptation

To address potential catastrophic forgetting during fine-tuning, we employ the following strategies: *Prompt Transfer Mechanism:* To ensure knowledge retention across domains, personalized prompts are directly transferred as $\hat{p}_i^o = p_i^o$, maintaining the original representations. Shared prompts undergo a transformation to align with the target domain's feature space, preserving adaptability. *Few-Shot Fine-tuning:* Transferred prompts are further refined using a small set of target domain samples (LT_t^{few}), enabling

Algorithm 3: PromptCD-E for NeuralCD.

```

1: Input: The cognitive diagnosis model  $\mathcal{M}$  based on the NeuralCD backbone.  $\mathcal{M}$ ,
   records  $\mathbf{LS}$  for pre-training and  $\mathbf{LT}_t^{few}$  in target domain  $t$  for fine-tuning.
2: Output: fine-tuned model  $\mathcal{M}$ , the transfer prompts  $\hat{\mathbf{p}}_{stu}^o$  and  $\hat{\mathbf{p}}_{sub}^d$ .
3: —Pre-training Stage—
4: while  $e_1 \leq Epoch_{Pretrain}$  do
5:   for  $\mathbf{LS}_s \in \{\mathbf{LS}_1, \mathbf{LS}_2, \dots, \mathbf{LS}_{|S|}\}$  do
6:     Initialize the students embedding  $\alpha_s^{orig}$ , the exercises embedding  $\beta_s^{orig}$ ,
     prompts  $\mathbf{p}_{stu}^o$ ,  $\mathbf{p}_{sub}^d$  and  $\mathcal{M}$ ;
7:     Enhance the representation of student and exercise in Eq.(3) and Eq.(4);
     connect  $\mathbf{p}_{stu}^o$  to  $\alpha_s^{orig}$ , and  $\mathbf{p}_{sub}^d$  to  $\beta_s^{orig}$ , obtaining  $\alpha_s^{out}$  and  $\beta_s^{out}$  after
     mapping.
8:     Input  $\alpha_s^{out}$  and  $\beta_s^{out}$  into  $\mathcal{M}$ . Specifically, subtract  $\beta_s^{out}$  from  $\alpha_s^{out}$ , multiply
     by the exercise and knowledge vectors, and pass through NeuralCD to get
     the predicted score  $\mathbf{y}_s$ ;
9:     Calculate the loss using  $\mathbf{LS}_s$  to update the model;
10:   end for
11: end while
12: —Fine-Tuning Stage—
13: while  $e_2 \leq Epoch_{FineTune}$  do
14:   Initialize the students embedding  $\alpha_t^{orig}$ , the exercises embedding  $\beta_t^{orig}$ ;
15:   Obtain the transfer prompts  $\hat{\mathbf{p}}_{stu}^o$  and  $\hat{\mathbf{p}}_{sub}^d$  in Eq.(5) and Eq.(6);
16:   Activate improvement policy in Eq.(7);
17:   Enhance the representations in a manner similar to pre-training;
18:   Input  $\alpha_t^{out}$  and  $\beta_t^{out}$  to  $\mathcal{M}$  to obtain the final predicted score  $\mathbf{y}_t$ , as in the
   pre-training process;
19:   Calculate the loss using  $\mathbf{LT}_t^{few}$  to update the model.
20: end while

```

domain-specific adaptation while maintaining core knowledge integrity.

To evaluate the retention of source domain knowledge after fine-tuning, we selected the scenario where the target domain is math in the exercise-aspect CDCD task [Fig. 9(c)] and the scenario where the target domain is A-bin in the student-aspect CDCD task [Fig. 9(d)]. Based on this, we performed a reverse transfer of both personalized and shared prompts, adapted to the target domain, back to the source domain model. Specifically, MIRT-S represents the source domain model, MIRT-T denotes the target domain model that received transferred prompts from the source domain, and MIRT-T-S refers to the source domain model that received reverse-transferred prompts from the target domain. Results indicate that after adapting to the target domain, these prompts can still effectively retain key information from the source domain, without significant information loss or catastrophic forgetting. This demonstrates that during fine-tuning, the model can achieve a good balance between improving adaptability to the target domain and preserving knowledge from the source domain.

REFERENCES

- [1] Y. Liu, T. Zhang, X. Wang, G. Yu, and T. Li, "New development of cognitive diagnosis models," *FCS*, vol. 17, no. 1, 2023, Art. no. 171604.
- [2] F. Wang et al., "Neural cognitive diagnosis for intelligent education systems," in *Proc. AAAI*, 2020, pp. 6153–6161.
- [3] H. Ma et al., "Knowledge-sensed cognitive diagnosis for intelligent education platforms," in *Proc. CIKM*, 2022, pp. 1451–1460.
- [4] F. Liu, X. Hu, S. Liu, C. Bu, and L. Wu, "Meta multi-agent exercise recommendation: A game application perspective," in *Proc. SIGKDD*, New York, NY, USA, 2023, pp. 1441–1452.
- [5] Z. Huang et al., "Exploring multi-objective exercise recommendations in online education systems," in *Proc. CIKM*, 2019, pp. 1261–1270.
- [6] S. Wu, J. Wang, and W. Zhang, "Contrastive personalized exercise recommendation with reinforcement learning," in *Proc. IEEE TLT*, 2023.
- [7] J. Fan, Y. Jiang, Y. Liu, and Y. Zhou, "Interpretable MOOC recommendation: a multi-attention network for personalized learning behavior analysis," *Internet Res.*, vol. 32, no. 2, pp. 588–605, 2022.
- [8] H. Bi et al., "Quality meets diversity: A model-agnostic framework for computerized adaptive testing," in *Proc. ICDM*, Piscataway, NJ, USA: IEEE Press, 2020, pp. 42–51.
- [9] H.-S. Chang, H.-J. Hsu, and K.-T. Chen, "Modeling exercise relationships in e-learning: A unified approach," in *Proc. EDM*, 2015, pp. 532–535.
- [10] Z. Wang et al., "Instructions and guide for diagnostic questions: The NeurIPS 2020 Education Challenge," 2021, *arXiv:2007.12061*.
- [11] Q. Liu et al., "Survey of computerized adaptive testing: a machine learning perspective," 2024, *arXiv:2404.00712*.
- [12] F. Wang et al., "A survey of models for cognitive diagnosis: new developments and future directions," 2024, *arXiv:2407.05458*.
- [13] Y. Zhou et al., "Modeling context-aware features for cognitive diagnosis in student learning," in *Proc. SIGKDD*, 2021, pp. 2420–2428.
- [14] W. Gao et al., "RCD: Relation map driven cognitive diagnosis for intelligent education systems," in *Proc. SIGIR*, 2021, pp. 501–510.
- [15] F. Wang et al., "NeuralCD: A general framework for cognitive diagnosis," *IEEE Trans. Knowl. Data Eng.*, vol. 35, no. 8, pp. 8312–8327, Aug. 2022.
- [16] X. Wang, C. Huang, J. Cai, and L. Chen, "Using knowledge concept aggregation towards accurate cognitive diagnosis," in *Proc. CIKM*, 2021, pp. 2010–2019.
- [17] J. Li et al., "HierCDF: A Bayesian network-based hierarchical cognitive diagnosis framework," in *Proc. SIGKDD*, 2022, pp. 904–913.
- [18] S. Tong et al., "Incremental cognitive diagnosis for intelligent education," in *Proc. SIGKDD*, 2022, pp. 1760–1770.
- [19] M. D. Reckase and M. D. Reckase, *Multidimensional Item Response Theory Models*, New York, NY, USA: Springer, 2009.
- [20] Y. Lu et al., "SLP: A multi-dimensional and consecutive dataset from k-12 education," in *Proc. ICCE*, vol. 1, 2021, pp. 261–266.
- [21] W. Gao et al., "Leveraging transferable knowledge concept graph embedding for cold-start cognitive diagnosis," in *Proc. SIGIR*, 2023, pp. 983–992.
- [22] W. Gao et al., "Zero-1-to-3: Domain-level zero-shot cognitive diagnosis via one batch of early-bird students towards three diagnostic objectives," in *Proc. AAAI*, 2024.
- [23] P. Liu, W. Yuan, J. Fu, Z. Jiang, H. Hayashi, and G. Neubig, "Pre-train, prompt, and predict: A systematic survey of prompting methods in natural language processing," *ACM Comput. Surv.*, vol. 55, no. 9, pp. 1–35, 2023.
- [24] C. Ge et al., "Domain adaptation via prompt learning," *IEEE TNNLS*, 2023, *arXiv:2202.06687*.
- [25] J. Liu et al., "Prompt learning with cross-modal feature alignment for visual domain adaptation," in *Proc. CAAI InterNat. Conf. Artif. Intell.*, Cham, Switzerland: Springer, 2022, pp. 416–428.
- [26] H. Wu and X. Shi, "Adversarial soft prompt tuning for cross-domain sentiment analysis," in *Proc. ACL*, 2022, pp. 2438–2447.
- [27] W. Zhao, A. Gupta, T. Chung, and J. Huang, "SPC: Soft prompt construction for cross domain generalization," in *Proc. Workshop RepL4NLP*, 2023, pp. 118–130.
- [28] L. Zhao et al., "ADPL: Adversarial prompt-based domain adaptation for dialogue summarization with knowledge disentanglement," in *Proc. SIGIR*, 2022, pp. 245–255.
- [29] S. E. Embretson and S. P. Reise, *Item Response Theory*, New York, NY, USA: Psychology Press, 2013.
- [30] L. Hu et al., "PTADISC: a cross-course dataset supporting personalized learning in cold-start scenarios," in *Proc. NeurIPS (NIPS)*, Red Hook, NY, USA: Curran Associates Inc., 2024.
- [31] J. De La Torre, "Dina model and parameter estimation: A didactic," *J. Educ. Behav. Statist.*, vol. 34, no. 1, pp. 115–130, 2009.
- [32] H. Qian, S. Liu, M. Li, B. Li, Z. Liu, and A. Zhou, "ORCDF: An oversmoothing-resistant cognitive diagnosis framework for student learning in online education systems," in *Proc. SIGKDD*, 2024, pp. 2455–2466.
- [33] J. Shen, H. Qian, W. Zhang, and A. Zhou, "Symbolic cognitive diagnosis via hybrid optimization for intelligent education systems," in *Proc. AAAI/AAAI/EAAI*, Palo Alto, CA, USA: AAAI Press, 2024.
- [34] M. Li, H. Qian, J. Lv, M. He, W. Zhang, and A. Zhou, "Foundation model enhanced derivative-free cognitive diagnosis," *FCS*, vol. 19, no. 1, 2025, Art. no. 191318.
- [35] S. Liu, Z. Zhou, Y. Liu, J. Zhang, and H. Qian, "Language representation favored zero-shot cross-domain cognitive diagnosis," in *Proc. SIGKDD*, 2025, pp. 1–8.

- [36] F. Petroni et al., "Language models as knowledge bases?" 2019, *arXiv:1909.01066*.
- [37] Z. Zhang and B. Wang, "Prompt learning for news recommendation," 2023, *arXiv:2304.05263*.
- [38] J. Li, W. Zhang, T. Wang, G. Xiong, A. Lu, and G. Medioni, "GPT4REC: A generative framework for personalized recommendation and user interests interpretation," 2023, *arXiv:2304.03879*.
- [39] Y. Hou et al., "Large language models are zero-shot rankers for recommender systems," 2023, *arXiv:2305.08845*.
- [40] X. Wu, H. Zhou, W. Yao, X. Huang, and N. Liu, "Towards personalized cold-start recommendation with prompts," 2023, *arXiv:2306.17256*.
- [41] L. Li, Y. Zhang, and L. Chen, "Personalized prompt learning for explainable recommendation," *ACM TOIS*, vol. 41, no. 4, pp. 1–26, 2023.
- [42] Y. Gu, X. Han, Z. Liu, and M. Huang, "PPT: pre-trained prompt tuning for few-shot learning," 2021, *arXiv:2109.04332*.
- [43] Y. Zhou et al., "Large language models are human-level prompt engineers," 2022, *arXiv:2211.01910*.
- [44] V. Liu and L. B. Chilton, "Design guidelines for prompt engineering text-to-image generative models," in *Proc. CHI*, 2022, pp. 1–23.
- [45] V. Sanh et al., "Multitask prompted training enables zero-shot task generalization," 2021, *arXiv:2110.08207*.