

Temporal Smoothness Doubly Robust Learning for Debiased Knowledge Tracing

Peilin Zhan, Wei Chen, Weilin Chen, Shuyi Pan and Ruichu Cai*

School of Computer Science and Technology, Guangdong University of Technology, Guangzhou, China
 zhanpl01@outlook.com, {chenweidelight, chenweilin.chn, pansy0174, cairuichu}@gmail.com

Abstract

Knowledge Tracing (KT) is fundamental to intelligent education systems, yet relies on educational logs that are selectively observed. The non-random nature of exercise recommendations and student choices inevitably induces severe selection bias. Most existing KT methods neglect this issue, training on observed logs using standard empirical risk, which yields biased mastery estimates and accumulates errors in subsequent recommendations. To address this, we introduce a doubly robust (DR) formulation for KT that integrates a propensity model with an error imputation model, theoretically guaranteeing unbiasedness if either model is accurate. Beyond unbiasedness, in the sequential setting of KT, we identify that the estimator’s performance is compromised by variance-dependent stochastic deviations that accumulate over time, thereby causing training instability and limiting performance. To mitigate this, we derive a generalization bound that explicitly characterizes the impact of estimator variance and identifies temporal smoothness as a key factor in controlling it. Building on these theoretical insights, we propose the Temporal Smoothness Doubly Robust (TSDR) framework. TSDR jointly optimizes the KT predictor and the imputation model with a smoothness regularizer, effectively reducing variance while preserving the unbiasedness guarantee of DR. Experiments on multiple real-world benchmarks demonstrate that TSDR consistently enhances various state-of-the-art KT backbones, underscoring the vital role of principled bias correction in KT.

1 Introduction

In intelligent education, tracing what students know is of great importance for adaptive learning [Anderson *et al.*, 1985; Villegas-Ch *et al.*, 2025]. Knowledge Tracing (KT) serves this role by modeling students’ dynamic knowledge states based on their historical learning interactions, aiming to predict how well they will master future questions.

*Corresponding author.

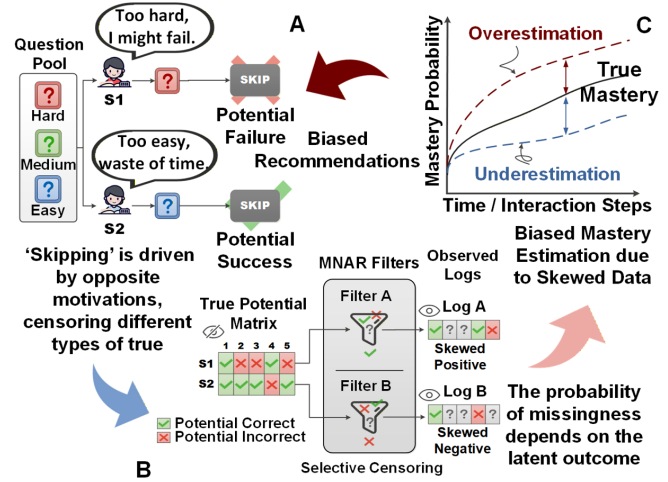


Figure 1: The causal mechanism of selection bias in KT: student skipping acts as a bias filter that systematically distorts the observed learning traces, causing models to misinterpret missingness and gradually internalize these distortions, ultimately learning biased knowledge states.

However, despite the rapid progress of deep learning-based KT models [Piech *et al.*, 2015; Pandey and Karypis, 2019; Ghosh *et al.*, 2020; Huang *et al.*, 2023; Cui *et al.*, 2025], their predictions often exhibit systematic deviation from students’ actual knowledge. This discrepancy raises a fundamental question: are current KT models truly tracing students’ knowledge, or just fitting their observed behaviors?

The reason for this discrepancy lies in the nature of the data on which current KT models are trained. The recorded data exhibit bias since whether student-question interactions are observed may depend on the students’ ability and the questions’ difficulty, a phenomenon known as the Missing Not At Random (MNAR) property. Taking Fig. 1 as an example, in online learning scenarios, students often strategically skip questions they deem “too difficult” or “boring” [Liu *et al.*, 2025]. This behavior causes a lack of specific question types in the data. Ignoring this missingness mechanism leads standard KT models to confound genuine knowledge states with interaction patterns, thereby introducing severe selection bias. This results in skewed proficiency estimates and diminished generalization capabilities. The issue is further compounded

in adaptive systems, since recommendations are tailored to the student’s estimated level via the Zone of Proximal Development (ZPD) [Vygotsky *et al.*, 1978]: biased estimates trigger biased recommendations, creating a feedback loop that exacerbates the bias [Carless, 2019].

Addressing such a selection bias necessitates disentangling a student’s actual performance from the observation probability. This requires modeling two factors: which interactions are logged and what the responses would have been for unlogged interactions. Accordingly, we adopt a Doubly Robust (DR) risk estimator that integrates selection correction via Inverse Propensity Weighting (IPW) [Schnabel *et al.*, 2016; Thomas and Brunskill, 2016] and potential performance estimation via Error Imputation (EIB) [Steck, 2013]. This construction enjoys the double robustness property: the resulting risk estimate is unbiased as long as either the propensity model or the imputation model is accurate. Nevertheless, directly applying DR in KT can be unstable: the training loss is prone to oscillation if imputation models are inaccurate [Wang *et al.*, 2019]. This issue is particularly critical in educational logs where responses are inherently noisy due to guessing and slipping [Ma *et al.*, 2024; Baker *et al.*, 2008b; Wise, 2017], and such noise triggers error propagation and accumulation throughout deep sequential models, leading to significant error amplification.

To bridge this gap, we propose a Temporal Smoothness Doubly Robust (TSDR) learning framework for knowledge tracing.¹ Rather than treating the instability as an inevitable side effect, we derive a generalization bound specifically for the KT context. This theoretical analysis reveals that the estimation risk is dominated by the instability of imputation errors. Guided by this insight, we design a joint learning paradigm where the KT predictor and the imputation model are co-optimized. Crucially, we introduce a temporal smoothness constraint to regularize the latent trajectory, which effectively tightens the generalization bound and suppresses variance, thereby preserving the unbiased nature of the DR estimator while ensuring training stability. Overall, the main contributions of this paper are summarized as follows:

- We tackle the selection bias in KT induced by the non-random nature of exercise recommendations and student engagement choices. By adopting a causal inference perspective, we model the data generation process and propose a DR estimator to obtain unbiased mastery estimates.
- We propose TSDR, a model-agnostic DR framework for bias correction that incorporates temporal smoothing to mitigate variance-related stochastic deviations. Supported by the proposed generalization bound, we employ a joint learning paradigm to co-optimize the KT and imputation models, ensuring robust performance.
- We conduct extensive experiments on 9 real-world benchmark datasets. The empirical results validate the efficacy of our framework, showing that TSDR improves the performance of various state-of-the-art KT baselines.

¹The extended version is available at <https://arxiv.org/pdf/2605.05958>

These findings highlight the practical value of addressing selection bias in knowledge tracing.

2 Related Work

2.1 Classic Knowledge Tracing Models

Early research in knowledge tracing is based on probabilistic graphical models, such as BKT [Corbett and Anderson, 1994]. BKT leverages Hidden Markov Models to characterize knowledge state transitions, and Item Response Theory (IRT) [Baker, 2001], which focuses on modeling item difficulty and student ability. After that, the advent of deep learning catalyzed a profound paradigm shift. DKT [Piech *et al.*, 2015] pioneered the use of Recurrent Neural Networks to capture the hidden dynamics of knowledge acquisition within student interaction sequences, demonstrating strong adaptability in handling the complexities of real-world scenarios. Subsequently, the research paradigm gravitated towards Transformer-based architectures [Vaswani *et al.*, 2017] for their superior ability to capture long-range dependencies. For instance, Self-Attentive Knowledge Tracing (SAKT) [Pandey and Karypis, 2019] first applied self-attention mechanisms to a student’s historical interaction sequence to dynamically weigh the relevance of past exercises. AKT [Ghosh *et al.*, 2020] further refined this by incorporating contextual information from the items themselves to mitigate interference from irrelevant exercises. More recent work has been dedicated to augmenting attention mechanisms to overcome their native inability to capture the dynamics of human forgetting or to robustly handle sequences of long lengths [Li *et al.*, 2024; Im *et al.*, 2023]. While their approaches differ, a significant body of work increasingly seeks to enhance model predictive performance from the perspective of learning high-quality representations for knowledge concepts or exercises. For instance, some studies have introduced Graph Neural Networks to model relationships among knowledge concepts [Nakagawa *et al.*, 2019], while others have considered relations between knowledge concepts and exercises [Mao *et al.*, 2025], among other complex relationships.

However, these methods largely overlook systematic selection bias from non-random assignments, limiting their ability to generalize to counterfactual scenarios.

2.2 Bias and Debiasing in Knowledge Tracing

A central challenge in the development of robust Knowledge Tracing (KT) models is the mitigation of various biases that can compromise their validity. Existing research has predominantly focused on two key areas: first, addressing the inherent noise present within student interaction data; and second, preventing the model from learning spurious correlations during the training process.

Debiasing from Noise in Student Responses. The first major line of research targets the inherent noise in student response data, which can distort the model’s inference of a student’s true knowledge state. This noise often stems from aberrant student behaviors. Early models explicitly incorporated parameters for guessing and slipping to account for such behaviors [Baker *et al.*, 2008b; Baker *et al.*, 2008a;

Zambrano *et al.*, 2024; Barrett *et al.*, 2024]. More recent work has explored nuanced modeling of factors like student disengagement [Gorgun and Bulut, 2022] and added small perturbations to input embeddings [Guo *et al.*, 2021] to improve robustness. State-of-the-art approaches now employ adversarial training and contrastive learning to generate debiased knowledge representations that are resilient to abnormal responses, thereby isolating the most reliable signals for prediction [Lv *et al.*, 2025; Ma *et al.*, 2024], also acknowledging that student actions do not always directly reflect their knowledge [Lv *et al.*, 2025]. While mitigating response noise, they focus on denoising rather than data recovery, failing to account for selection bias from non-random missingness.

Addressing Spurious Correlations and Causal Confusion.

The second major area of research addresses the risk of the model learning spurious correlations from confounding factors, which is often termed “causal confusion” [de Haan *et al.*, 2019]. This occurs when a model learns shortcuts from data artifacts instead of isolating the student’s true knowledge. These confounders can stem from various sources, including the model’s own architectural limitations. To counteract confounding more comprehensively, researchers have utilized a spectrum of techniques, many of which are predicated on the principles of causal inference. For instance, to address confounding from the imbalanced informativeness of historical responses, DR4KT introduces a reweighting framework [Cui *et al.*, 2025]. Other approaches use formal causal modeling to disentangle effects. The CORE framework mitigates answer bias by modeling and subtracting the direct causal effect of the question on the response [Cui *et al.*, 2023]. Similarly, DisKT uses a disentangled modeling approach and a causal subtraction mechanism to remove confounding from uneven performance distributions across concepts [Zhou *et al.*, 2025]. For more complex scenarios involving unobserved confounders, recent work has leveraged the front-door criterion from causal inference, implementing it through causal self-attention mechanisms to identify the true, unconfounded pathway from knowledge to performance [Huang *et al.*, 2024]. While recent causal-aware approaches attempt to mitigate specific confounders, they primarily focus on feature-level disentanglement rather than correcting the systemic selection bias in the data generation process.

2.3 MNAR and DR Debiasing

While existing KT approaches address biases regarding *what* is recorded and *how* it is modeled, they largely overlook a fundamental bias rooted in *why* interactions are recorded: selection bias. This challenge is closely related to the MNAR problem in Recommender Systems (RS), where doubly robust learning mitigates selection bias by combining propensity weighting with error imputation [Dudík *et al.*, 2011; Wang *et al.*, 2019]. Beyond RS, DR has also been explored in structured observational settings such as networked interference [Chen *et al.*, 2024]. Similar concerns arise in time-series imputation, where causal perspectives emphasize missing mechanisms in temporal data [Cai *et al.*, 2025]. However, KT differs from these settings by modeling evolving knowledge states rather than static preferences, networked effects,

or imputed values. Directly applying existing DR estimators to KT may therefore cause high-variance corrections and error accumulation along the learning sequence. We thus adapt DR to KT and introduce temporal smoothness to stabilize sequential training.

3 Problem Formulation

The objective of KT is to predict a student’s performance on an upcoming question by modeling their knowledge state from past interactions. Without loss of generality, we define a student’s historical interaction sequence as $Seq_t := \{(c_1, r_1), (c_2, r_2), \dots, (c_t, r_t)\}$, where c_t and r_t denote the concept from the concept set $\mathcal{C}$ and the binary response correctness ($r_t \in \{0, 1\}$) at time step t respectively. Given the next concept c_{t+1} , the KT model aims to first estimate the student’s latent knowledge state $\hat{h}_t = f_\theta(Seq_t)$, and then use this state to predict the probability of a correct answer, $\hat{r}_{t+1} = g_\phi(\hat{h}_t, c_{t+1})$. The model parameters θ and ϕ are typically optimized by minimizing the prediction error, e.g., binary cross-entropy (BCE), on the observed data.

Since interactions are selectively acquired (MNAR) rather than random based on recommendation policies or student preferences, minimizing error on observed data results in a biased estimator. We typically analyze this bias and derive a causal-aware solution to mitigate it, given the students’ historical interaction sequences data.

4 Theoretical Analysis

To address the selection bias identified in the problem formulation, we reframe the KT task from the perspective of counterfactual estimation [Pearl *et al.*, 2016] for model performance. We formalize the discrepancy between the ideal true risk and the naive empirical risk, and then derive the DR estimator as a principled solution.

4.1 Revisiting Estimation Problem

Let data $\mathcal{D}_{\text{full}} = \{(t, c) \mid 1 \leq t \leq T, c \in \mathcal{C}\}$ with T being the maximum length of the interaction sequence denote the ideal full set of potential interactions for a student sequence. We define the true risk $\mathcal{R}_{\text{true}}$ as the average prediction error over this full distribution of data:

$$\mathcal{R}_{\text{true}} = \mathbb{E}_{(t,c) \sim \mathcal{D}_{\text{full}}} [e_{t+1,c}] = \frac{1}{N} \sum_{t=1}^{T-1} \sum_{c=1}^{|\mathcal{C}|} e_{t+1,c}, \quad (1)$$

where $N = (T - 1) \times |\mathcal{C}|$ represents the total number of prediction targets, and $e_{t+1,c} = -(r_{t+1,c} \log(\hat{r}_{t+1,c}) + (1 - r_{t+1,c}) \log(1 - \hat{r}_{t+1,c}))$ is the BCE loss for skill c at time t .

In practice, we only observe a subset of interactions where the observation indicator $o_{t+1,c} = 1$. The **Naive Estimator**, which corresponds to the standard KT objective, calculates the risk solely on observed data:

$$\hat{\mathcal{R}}_{\text{naive}} = \frac{\sum o_{t+1,c} \cdot e_{t+1,c}}{\sum o_{t+1,c}}. \quad (2)$$

Since the observation probability, i.e., propensity $p_{t+1,c} = P(o_{t+1,c} = 1 | h_t, c)$, is non-uniform due to the MNAR mechanism, $\mathbb{E}[\hat{\mathcal{R}}_{\text{naive}}] \neq \mathcal{R}_{\text{true}}$. This estimation bias misguides

the optimization, causing the model to overfit to the specific selection policy rather than learning the true knowledge dynamics.

4.2 Constructing DR Estimator

To obtain an unbiased estimate of $\mathcal{R}_{\text{true}}$, we combine two distinct approaches: the imputation-based method modeling $\hat{e}_{t+1,c}$ and the propensity-based method modeling $\hat{p}_{t+1,c}$.

We introduce the DR estimator, denoted as $\hat{\mathcal{R}}_{\text{DR}}$. Instead of viewing it merely as a loss function, we define it as a statistical estimator for the unobserved true risk:

$$\hat{\mathcal{R}}_{\text{DR}} = \frac{1}{N} \sum_{t=1}^{T-1} \sum_{c=1}^{|C|} \left[\underbrace{\hat{e}_{t+1,c}}_{\text{Imputation}} + \underbrace{\frac{o_{t+1,c}}{\hat{p}_{t+1,c}}(e_{t+1,c} - \hat{e}_{t+1,c})}_{\text{Correction Term}} \right], \quad (3)$$

where N represents the total number of entries, T is the maximum sequence length, and C denotes the set of knowledge concepts. The intuition behind this formulation is to leverage the imputed error $\hat{e}_{t+1,c}$ as a **baseline estimate** for the risk on the full distribution. Since the imputation model alone may be biased, the second term, $\frac{o_{t+1,c}}{\hat{p}_{t+1,c}}(e_{t+1,c} - \hat{e}_{t+1,c})$, acts as a **bias correction**. It adjusts the baseline using the residual between the true error $e_{t+1,c}$ and the imputed error $\hat{e}_{t+1,c}$ for observed interactions, weighted by the inverse propensity. This structure explicitly demonstrates the “double robustness” property: if the imputation is accurate (i.e., $\hat{e} \approx e$), the correction term vanishes, effectively minimizing variance; conversely, if the propensity model is accurate, the correction term ensures the estimator remains unbiased even if the baseline imputation is poor.

4.3 Theoretical Guarantee from Bias Analysis

To show why this estimator is preferred, we analyze its bias with respect to the true risk. The bias is defined as the difference between the expected value of the estimator over the observation distribution O and the true target. The complete theoretical framework and main results are presented in Appendix E, with detailed proofs provided in Appendix F.

For clearly, we define the imputation error as $\delta_{t,c} := e_{t,c} - \hat{e}_{t,c}$ and the propensity error as $\Delta_{t,c} := 1 - \frac{p_{t,c}}{\hat{p}_{t,c}}$.

Lemma 1 (Bias of the DR Estimator). *The bias of the DR estimator is bounded by the product of the errors of the two auxiliary models:*

$$\left| \mathbb{E}_O[\hat{\mathcal{R}}_{\text{DR}}] - \mathcal{R}_{\text{true}} \right| \leq \frac{1}{N} \sum_{(t,c) \in \mathcal{D}_{\text{full}}} |\Delta_{t+1,c} \cdot \delta_{t+1,c}|. \quad (4)$$

Proof. (Sketch) By linearity of expectation, $\mathbb{E}_O[o_{t,c}] = p_{t,c}$. Substituting this into the expectation of $\hat{\mathcal{R}}_{\text{DR}}$:

$$\begin{aligned} \mathbb{E}_O[\hat{\mathcal{R}}_{\text{DR}}] &= \frac{1}{N} \sum_{t,c} \left(\hat{e}_{t+1,c} + \frac{p_{t+1,c}}{\hat{p}_{t+1,c}}(e_{t+1,c} - \hat{e}_{t+1,c}) \right) \\ &= \frac{1}{N} \sum_{t,c} \left(e_{t+1,c} \frac{p_{t+1,c}}{\hat{p}_{t+1,c}} + \hat{e}_{t+1,c} \left(1 - \frac{p_{t+1,c}}{\hat{p}_{t+1,c}} \right) \right). \end{aligned} \quad (5)$$

Subtracting $\mathcal{R}_{\text{true}} = \frac{1}{N} \sum e_{t+1,c}$ from this expectation yields the bias term proportional to $\delta_{t+1,c} \cdot \Delta_{t+1,c}$. $\square$

Lemma 1 confirms the **Double Robustness** property: the estimation is unbiased if **either** the propensity model or the imputation model is accurate. More attractively, the bias depends on a **product** of the two model errors. When both models are learned with sufficiently expressive estimators (e.g., neural networks) and their estimation errors can be controlled, the DR estimator achieves a **second-order error decay**, resulting in an error that is smaller than that of estimators typically dominated by a single estimation error.

4.4 Generalization Bound and Variance Control

While the DR estimator theoretically guarantees unbiasedness under mild conditions, its practical application on finite interaction sequences is often hindered by stochastic fluctuations. To rigorously quantify this risk, we derive a generalization bound that accounts for both the bias and the concentration properties of the estimator.

Theorem 1 (Generalization Bound). *For any finite hypothesis space $\mathcal{H}$, with probability at least $1 - \eta$, the true risk of the optimal prediction is bounded by:*

$$\begin{aligned} \mathcal{R}_{\text{true}} \leq & \underbrace{\hat{\mathcal{R}}_{\text{DR}}}_{\text{Empirical Risk}} + \underbrace{\frac{1}{N} \sum_{(t,c) \in \mathcal{D}_{\text{full}}} |\Delta_{t+1,c} \delta_{t+1,c}|}_{\text{Bias Term}} \\ & + \underbrace{\sqrt{\frac{\ln(2|\mathcal{H}|/\eta)}{2N^2} \sum_{(t,c) \in \mathcal{D}_{\text{full}}} \left(\frac{\delta_{t+1,c}}{\hat{p}_{t+1,c}} \right)^2}}_{\text{Variance Term}}. \end{aligned} \quad (6)$$

Proof. (Sketch) The bound is derived by decomposing the risk into bias and variance components. The bias term follows from Lemma 1. The variance term is bounded using the Azuma-Hoeffding inequality [Azuma, 1967], accounting for the sequential nature of student interactions, combined with a union bound over $\mathcal{H}$. $\square$

Theorem 1 reveals a critical insight: the total generalization error is dominated not just by the bias product but also by the Variance Term, which scales with the squared imputation error $\delta_{t,c}^2$. Inaccuracy in the imputation model results in large error deviations $\delta_{u,i}$, which loosens the generalization bound and in turn leads to unstable training. To minimize $\delta_{t,c}$ and tighten the generalization bound, we introduce a temporal smoothness constraint (Section 5.1). This constraint works by regularizing the latent trajectory. Furthermore, we adopt a joint learning approach to couple the optimization processes (Section 5.2).

5 Method

We re-conceptualize the challenge of debiasing a temporal learning process, treating the knowledge state at each distinct time step as a temporal “user” and operating under the principle that each momentary state possesses its own distinct learning preferences and biases. Adopting this perspective allows us to apply a DR strategy to learn from these state-specific interactions in a robust manner.

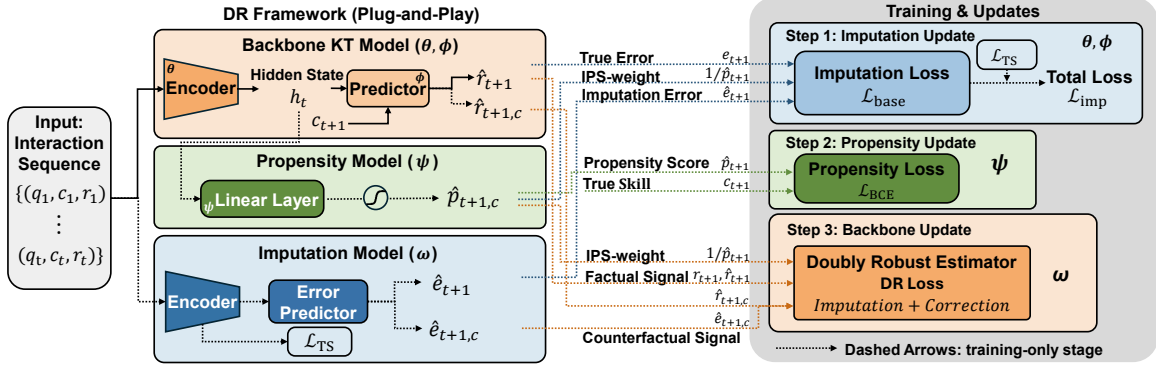


Figure 2: A Model-Agnostic DR Framework Overview.

5.1 Model Architecture

We propose a **Temporal-Smooth Doubly Robust (TSDR)** learning framework that jointly learns a propensity model, an error imputation model, and a KT model. The overall architecture of our framework is presented in Figure 2.

Propensity Model for Interaction Selection

We learn the propensity score to model the interaction probability as:

$$\hat{p}_{t+1,c} = P(o_{t+1,c} = 1 \mid \hat{h}_t, c). \quad (7)$$

We model this as a **Multivariate Bernoulli distribution** for all concepts using a sigmoid output layer. The model is trained via BCE loss, treating unobserved concepts as implicit negatives.

Error Imputation Model for Counterfactual Prediction

The second component is the error imputation model, parameterized by ω . Unlike the prediction model that focuses on knowledge mastery, this module aims to capture the dynamics of prediction residuals. To decouple these distinct objectives, we employ a separate encoder for the imputation task. Let h_t^{imp} denote the imputation-specific latent state derived from the interaction sequence Seq_t :

$$h_t^{imp} = \text{Enc}_\omega(Seq_t), \quad (8)$$

where Enc_ω denotes the internal knowledge state extractor of the imputation model, similar to the KT backbone. The imputed error $\hat{e}_{t+1}$ is then predicted based on this state and a potential next concept:

$$\hat{e}_{t+1,c} = g_\omega(h_t^{imp}, c). \quad (9)$$

Optimization Objective. The training of the imputation model consists of two components. First, the model aims to regress the true prediction errors e_{t+1} based on the observed data. To address the distribution shift, we adopt an Inverse Propensity Score (IPS) weighted Mean Squared Error (MSE) as the base objective:

$$\mathcal{L}_{\text{base}}(\omega) = \frac{1}{T-1} \sum_{t=1}^{T-1} \frac{(\hat{e}_{t+1} - e_{t+1})^2}{\hat{p}_{t+1}}, \quad (10)$$

where we focus solely on observation loss. Then, guided by the theoretical analysis in Theorem 1, we incorporate a

Temporal Smoothness constraint on the latent states h_t^{imp} . This aligns with the Power Law of Practice [Newell and Rosenbloom, 2013], which characterizes skill acquisition as a smooth, continuous power-function trajectory rather than a series of abrupt changes. By explicitly restricting the path length of H^{imp} , this constraint reduces the Rademacher complexity and tightens the generalization bound. The total loss function is formulated as:

$$\mathcal{L}_{\text{imp}}(\omega) = \mathcal{L}_{\text{base}}(\omega) + \frac{\lambda}{T-1} \sum_{t=1}^{T-1} (h_t^{imp} - h_{t-1}^{imp})^2, \quad (11)$$

where λ is the coefficient controlling the strength of the smoothness regularization.

KT Model with Doubly Robust Estimator

As defined in our problem formulation, the core of our approach is the KT model, which consists of an encoder f_θ and a predictor g_ϕ . The encoder $f_\theta(Seq_t)$ processes the interaction history to produce the knowledge state $\hat{h}_t$. The predictor $g_\phi(\hat{h}_t, c_{t+1})$ then outputs the probability of a correct answer $\hat{r}_{t+1,c}$. Our ultimate goal is to learn the parameters θ and ϕ in an unbiased manner. The vanilla prediction loss for a single step is given by the BCE:

$$e_{t+1} = -r_{t+1} \log(\hat{r}_{t+1}) - (1 - r_{t+1}) \log(1 - \hat{r}_{t+1}). \quad (12)$$

Doubly Robust Estimator. Adhering to the principle of double robustness, we train the KT model by minimizing the doubly robust loss instead of the naive BCE loss. As restated in Eq. (3), the loss is given by:

$$\mathcal{L}_{\text{DR}}(\theta, \phi) = \frac{1}{N} \sum_{t=1}^{T-1} \sum_{c \in \mathcal{C}} (\hat{e}_{t+1,c} + w_{t+1,c} \Delta e_{t+1,c}), \quad (13)$$

where $w_{t+1,c} = o_{t+1,c} / \hat{p}_{t+1,c}$ and $\Delta e_{t+1,c} = e_{t+1,c} - \hat{e}_{t+1,c}$.

5.2 Joint Learning Approach

To mitigate the impact of imputation errors on prediction, we derive a generalization bound (Appendix F) following [Wang et al., 2019]. Motivated by this theoretical insight, we co-optimize the parameters $(\theta, \phi, \psi, \omega)$ via an alternating strategy, as detailed in Algorithm 1. This joint procedure allows the models to mutually regularize each other, leading to a more robust estimation of the true prediction inaccuracy and ultimately, a more effective and unbiased KT model.

Algorithm 1 Joint Learning with TSDR

Input: Observed interaction data $\mathcal{D}_{obs}$, Concept set $\mathcal{C}$.

Parameters: Learning rates α .

Output: Trained KT model parameters θ, ϕ

```

1: Initialize parameters  $\theta, \phi, \psi, \omega$ .
2: while stopping criteria not met do
3:   Sample a mini-batch  $\mathcal{B}$  from  $\mathcal{D}_{obs}$ .
4:   /* Phase 1: Update Auxiliary Models */
5:   Update imputation model:  $\omega \leftarrow \omega - \alpha \nabla_{\omega} \mathcal{L}_{imp}(\omega)$ 
6:   Update propensity model:  $\psi \leftarrow \psi - \alpha \nabla_{\psi} \mathcal{L}_{prop}(\psi)$ 
7:   /* Phase 2: Update the KT Model */
8:   Based on observed history in  $\mathcal{B}$ , compute counterfactual errors  $\hat{e}_{t+1,c}$  for all concepts  $c \in \mathcal{C}$ .
9:   Calculate  $\mathcal{L}_{DR}(\theta, \phi)$  by summing errors over  $\mathcal{C}$ .
10:   $\theta, \phi \leftarrow (\theta, \phi) - \alpha \nabla_{\theta, \phi} \mathcal{L}_{DR}(\theta, \phi)$ 
11: end while
12: return Optimized KT parameters  $\theta, \phi$ .

```

6 Experiment

In this section, our evaluation aims to answer the following research questions:

- **RQ1:** Does TSDR consistently enhance the performance of SOTA baselines on real-world datasets?
- **RQ2:** Can TSDR mitigate performance degradation under varying MNAR levels in controlled simulations?
- **RQ3:** How sensitive is the model to the intensity of the temporal smoothness constraint?

6.1 Experimental Setup

Benchmark Dataset. We evaluate the proposed framework on 9 publicly available benchmark datasets covering diverse domains: Spanish², ASSISTments17³, Slepemapy⁴, Algebra05⁵, [Prob, Linux, Database, Comp]⁶ and EdNet⁷. These datasets vary significantly in scale and density, offering a comprehensive testbed for robustness. Details are available in Appendix A.

Synthetic Dataset. To evaluate robustness against the bias, we generated 6 synthetic datasets using a controllable MNAR mechanism with bias degrees $\gamma \in [0, 0.999]$. Generation details and statistics are provided in Appendix B.

Baselines. To validate TSDR as a generic strategy, we apply it to 6 SOTA baselines (DKT [Piech *et al.*, 2015], AKT [Ghosh *et al.*, 2020], simpleKT [Liu *et al.*, 2023], FoLiBiKT [Im *et al.*, 2023], SparseKT [Huang *et al.*, 2023], and DisKT [Zhou *et al.*, 2025]) and compare them against their original versions. Details of baselines are provided in Appendix C.

Implementation Details. We adopt the pipeline from the work in [Zhou *et al.*, 2025] using student-stratified 5-fold cross-validation. A 10% validation set determines early stopping after 15 epochs of non-improvement. We truncate sequences to the latest 50 interactions and train models using Adam with a learning rate of 0.001, a batch size of 64, and an embedding dimension of 64. The random seed and dropout rate are set to 42 and 0.05, respectively. Furthermore, we tune the hyperparameter λ that controls loss strength within [0.1, 0.3, 0.5, 0.7, 1, 2]. All experiments run on an NVIDIA RTX 4090 GPU and are evaluated via AUC, ACC, and RMSE following [Shen *et al.*, 2021; Im *et al.*, 2023; Zhou *et al.*, 2025]. The source code and supplementary material are available at <https://github.com/plzhan/TSDR>.

6.2 Results

Empirical Study (RQ1)

Table 1 summarizes the performance of TSDR applied to six baseline models across nine datasets. Comparing the original performance with the TSDR-enhanced results, we have the following observations: 1) Regarding the primary evaluation metric of AUC, TSDR yields improvements across diverse model architectures. These gains are particularly notable on datasets characterized by high sparsity or strong student autonomy. For instance, TSDR improves AKT by 5.06% on the Prob dataset and SparseKT by 4.76% on Assist17. These results suggest that our framework helps recover mastery patterns by mitigating bias in MNAR data. 2) The magnitude of improvement appears to correlate with the data collection mechanism. While TSDR provides substantial gains on open-ended platforms, the improvements on EdNet and Linux are relatively limited, as exemplified by an AUC gain of approximately 1% on EdNet. This behavior is likely attributable to EdNet’s specific design, which enforces a mandatory “bundle” completion policy [Choi *et al.*, 2020]. Such structural constraints tend to limit selection bias, rendering the missing data mechanism closer to random compared to other datasets, thereby reducing the potential scope for bias correction. 3) Regarding secondary metrics such as ACC and RMSE, TSDR consistently improves calibration. Even on datasets where AUC gains plateau, such as Linux and Database, TSDR frequently reduces RMSE and improves ACC. A specific example is the 1.70% reduction for SparseKT on Linux. This indicates that counterfactual imputation may contribute to better-calibrated probability estimates, even when the ranking performance remains stable.

Experiments on Synthetic Data (RQ2)

To evaluate robustness, we tested TSDR on synthetic datasets with MNAR degrees (γ) ranging from 0.0 to 0.999. As shown in Figure 3, we observe varying degrees of performance degradation across all baselines as the bias severity increases, particularly visible in AUC. However, equipping SOTA models with TSDR consistently yields improvements across the entire spectrum. Specifically for RMSE, with the exception of DKT, TSDR is highly effective at reducing error when MNAR is low, although this advantage diminishes slightly beyond $\gamma > 0.8$. Since TSDR is a model-agnostic strategy, it relies on the backbone’s modeling capacity to maximize ef-

²<https://github.com/robert-lindsey/WCRP>

³<https://sites.google.com/view/assistentmentsdatamining/dataset>

⁴<https://www.fi.muni.cz/adaptivelearning/?a=data>

⁵<https://pslcdatashop.web.cmu.edu/KDDCup/>

⁶<https://github.com/wahr0411/PTADisc>

⁷<https://github.com/rriid/ednet>

Dataset	Metric	DKT	AKT	simpleKT	FoLiBiKT	SparseKT	DisKT
Spanish	AUC $\uparrow$	0.7798 / 0.8018 ($\uparrow$ 2.82%)	0.7967 / 0.8034 ($\uparrow$ 0.84%)	0.8098 / 0.8188 ($\uparrow$ 1.11%)	0.7992 / 0.8067 ($\uparrow$ 0.94%)	0.8019 / 0.8197 ($\uparrow$ 2.22%)	0.8243 / 0.8259 ($\uparrow$ 0.19%)
	ACC $\uparrow$	0.7250 / 0.7362 ($\uparrow$ 1.54%)	0.7386 / 0.7420 ($\uparrow$ 0.46%)	0.7405 / 0.7494 ($\uparrow$ 1.20%)	0.7332 / 0.7388 ($\uparrow$ 0.76%)	0.7324 / 0.7510 ($\uparrow$ 2.54%)	0.7520 / 0.7582 ($\uparrow$ 0.82%)
	RMSE $\downarrow$	0.4301 / 0.4195 ($\downarrow$ 2.46%)	0.4294 / 0.4241 ($\downarrow$ 1.23%)	0.4214 / 0.4186 ($\downarrow$ 0.66%)	0.4292 / 0.4264 ($\downarrow$ 0.65%)	0.4319 / 0.4204 ($\downarrow$ 2.66%)	0.4199 / 0.4096 ($\downarrow$ 2.45%)
ASSISTments17	AUC	0.6358 / 0.6472 ($\uparrow$ 1.79%)	0.6843 / 0.7012 ($\uparrow$ 2.47%)	0.6692 / 0.6905 ($\uparrow$ 3.18%)	0.6826 / 0.6994 ($\uparrow$ 2.46%)	0.6800 / 0.7124 ($\uparrow$ 4.76%)	0.7223 / 0.7324 ($\uparrow$ 1.40%)
	ACC	0.6172 / 0.6249 ($\uparrow$ 1.25%)	0.6525 / 0.6617 ($\uparrow$ 1.41%)	0.6430 / 0.6590 ($\uparrow$ 2.49%)	0.6538 / 0.6619 ($\uparrow$ 1.24%)	0.6442 / 0.6693 ($\uparrow$ 3.90%)	0.6621 / 0.6784 ($\uparrow$ 2.46%)
	RMSE	0.4807 / 0.4774 ($\downarrow$ 0.69%)	0.4762 / 0.4651 ($\downarrow$ 2.33%)	0.4791 / 0.4662 ($\downarrow$ 2.28%)	0.4750 / 0.4659 ($\downarrow$ 1.92%)	0.4767 / 0.4624 ($\downarrow$ 3.00%)	0.4740 / 0.4641 ($\downarrow$ 2.09%)
Slepemapy	AUC	0.6806 / 0.6937 ($\uparrow$ 1.92%)	0.7127 / 0.7208 ($\uparrow$ 1.14%)	0.7108 / 0.7143 ($\uparrow$ 0.49%)	0.7120 / 0.7199 ($\uparrow$ 1.11%)	0.7478 / 0.7758 ($\uparrow$ 3.74%)	0.7283 / 0.7345 ($\uparrow$ 0.85%)
	ACC	0.7740 / 0.7795 ($\uparrow$ 0.71%)	0.7796 / 0.7833 ($\uparrow$ 0.47%)	0.7743 / 0.7830 ($\uparrow$ 1.12%)	0.7799 / 0.7843 ($\uparrow$ 0.56%)	0.7845 / 0.7944 ($\uparrow$ 1.26%)	0.7805 / 0.7861 ($\uparrow$ 0.72%)
	RMSE	0.4021 / 0.3977 ($\downarrow$ 1.09%)	0.3948 / 0.3915 ($\downarrow$ 0.84%)	0.3969 / 0.3927 ($\downarrow$ 1.06%)	0.3962 / 0.3915 ($\downarrow$ 1.19%)	0.3866 / 0.3768 ($\downarrow$ 2.53%)	0.3901 / 0.3886 ($\downarrow$ 0.38%)
Algebra05	AUC	0.7538 / 0.7669 ($\uparrow$ 1.74%)	0.7698 / 0.7797 ($\uparrow$ 1.29%)	0.7677 / 0.7809 ($\uparrow$ 1.71%)	0.7718 / 0.7781 ($\uparrow$ 0.82%)	0.7677 / 0.7772 ($\uparrow$ 1.24%)	0.7769 / 0.7862 ($\uparrow$ 1.20%)
	ACC	0.7748 / 0.7852 ($\uparrow$ 1.34%)	0.7794 / 0.7825 ($\uparrow$ 0.40%)	0.7734 / 0.7834 ($\uparrow$ 1.29%)	0.7813 / 0.7857 ($\uparrow$ 0.56%)	0.7694 / 0.7876 ($\uparrow$ 2.37%)	0.7817 / 0.7912 ($\uparrow$ 1.22%)
	RMSE	0.4013 / 0.3895 ($\downarrow$ 2.94%)	0.3955 / 0.3937 ($\downarrow$ 0.46%)	0.3980 / 0.3901 ($\downarrow$ 1.98%)	0.3911 / 0.3904 ($\downarrow$ 0.18%)	0.3954 / 0.3891 ($\downarrow$ 1.59%)	0.3958 / 0.3886 ($\downarrow$ 1.82%)
Prob	AUC	0.7278 / 0.7319 ($\uparrow$ 0.56%)	0.7347 / 0.7719 ($\uparrow$ 5.06%)	0.7447 / 0.7716 ($\uparrow$ 3.61%)	0.7409 / 0.7703 ($\uparrow$ 3.97%)	0.7597 / 0.7775 ($\uparrow$ 2.34%)	0.7617 / 0.7870 ($\uparrow$ 3.32%)
	ACC	0.6770 / 0.6822 ($\uparrow$ 0.77%)	0.6916 / 0.7136 ($\uparrow$ 3.18%)	0.6956 / 0.7110 ($\uparrow$ 2.21%)	0.6956 / 0.7126 ($\uparrow$ 2.44%)	0.7062 / 0.7197 ($\uparrow$ 1.91%)	0.7023 / 0.7262 ($\uparrow$ 3.40%)
	RMSE	0.4527 / 0.4510 ($\downarrow$ 0.38%)	0.4527 / 0.4407 ($\downarrow$ 2.65%)	0.4556 / 0.4466 ($\downarrow$ 1.98%)	0.4501 / 0.4403 ($\downarrow$ 2.18%)	0.4475 / 0.4411 ($\downarrow$ 1.46%)	0.4460 / 0.4364 ($\downarrow$ 2.15%)
Linux	AUC	0.7477 / 0.7521 ($\uparrow$ 0.59%)	0.8164 / 0.8162 ($\downarrow$ 0.02%)	0.8151 / 0.8162 ($\uparrow$ 0.13%)	0.8162 / 0.8163 ($\uparrow$ 0.01%)	0.8315 / 0.8413 ($\uparrow$ 1.18%)	0.8269 / 0.8278 ($\uparrow$ 0.11%)
	ACC	0.7690 / 0.7708 ($\uparrow$ 0.23%)	0.7975 / 0.7979 ($\uparrow$ 0.05%)	0.7963 / 0.7980 ($\uparrow$ 0.21%)	0.7970 / 0.7976 ($\uparrow$ 0.08%)	0.8039 / 0.8084 ($\uparrow$ 0.56%)	0.8045 / 0.8070 ($\uparrow$ 0.31%)
	RMSE	0.4009 / 0.3996 ($\downarrow$ 0.32%)	0.3760 / 0.3750 ($\downarrow$ 0.27%)	0.3763 / 0.3751 ($\downarrow$ 0.32%)	0.3759 / 0.3750 ($\downarrow$ 0.24%)	0.3702 / 0.3639 ($\downarrow$ 1.70%)	0.3705 / 0.3692 ($\downarrow$ 0.35%)
Database	AUC	0.7467 / 0.7493 ($\uparrow$ 0.35%)	0.8064 / 0.8079 ($\uparrow$ 0.19%)	0.8060 / 0.8067 ($\uparrow$ 0.09%)	0.8067 / 0.8071 ($\uparrow$ 0.05%)	0.8359 / 0.8422 ($\uparrow$ 0.75%)	0.8133 / 0.8156 ($\uparrow$ 0.28%)
	ACC	0.8294 / 0.8285 ($\downarrow$ 0.11%)	0.8383 / 0.8433 ($\uparrow$ 0.58%)	0.8395 / 0.8435 ($\uparrow$ 0.48%)	0.8392 / 0.8438 ($\uparrow$ 0.55%)	0.8478 / 0.8498 ($\uparrow$ 0.24%)	0.8460 / 0.8483 ($\uparrow$ 0.27%)
	RMSE	0.3562 / 0.3561 ($\downarrow$ 0.03%)	0.3415 / 0.3380 ($\downarrow$ 1.02%)	0.3408 / 0.3382 ($\downarrow$ 0.76%)	0.3410 / 0.3381 ($\downarrow$ 0.85%)	0.3302 / 0.3274 ($\downarrow$ 0.85%)	0.3360 / 0.3329 ($\downarrow$ 0.51%)
Comp	AUC	0.7073 / 0.7157 ($\uparrow$ 1.19%)	0.7912 / 0.7958 ($\uparrow$ 0.58%)	0.7929 / 0.7958 ($\uparrow$ 0.37%)	0.7914 / 0.7960 ($\uparrow$ 0.58%)	0.8142 / 0.8263 ($\uparrow$ 1.49%)	0.7986 / 0.8037 ($\uparrow$ 0.64%)
	ACC	0.8056 / 0.8077 ($\uparrow$ 0.26%)	0.8201 / 0.8232 ($\uparrow$ 0.38%)	0.8195 / 0.8243 ($\uparrow$ 0.59%)	0.8207 / 0.8245 ($\uparrow$ 0.46%)	0.8266 / 0.8314 ($\uparrow$ 0.58%)	0.8111 / 0.8298 ($\uparrow$ 2.31%)
	RMSE	0.3794 / 0.3773 ($\downarrow$ 0.55%)	0.3591 / 0.3558 ($\downarrow$ 0.92%)	0.3585 / 0.3558 ($\downarrow$ 0.75%)	0.3587 / 0.3552 ($\downarrow$ 0.98%)	0.3520 / 0.3454 ($\downarrow$ 1.87%)	0.3656 / 0.3529 ($\downarrow$ 3.47%)
EdNet	AUC	0.6602 / 0.6678 ($\uparrow$ 1.15%)	0.6868 / 0.6936 ($\uparrow$ 0.99%)	0.6924 / 0.6977 ($\uparrow$ 0.77%)	0.6885 / 0.6934 ($\uparrow$ 0.71%)	0.7282 / 0.7365 ($\uparrow$ 1.14%)	0.7273 / 0.7239 ($\downarrow$ 0.47%)
	ACC	0.6189 / 0.6270 ($\uparrow$ 1.31%)	0.6288 / 0.6351 ($\uparrow$ 1.00%)	0.6324 / 0.6451 ($\uparrow$ 2.01%)	0.6337 / 0.6402 ($\uparrow$ 1.03%)	0.6698 / 0.6800 ($\uparrow$ 1.52%)	0.6614 / 0.6698 ($\uparrow$ 1.27%)
	RMSE	0.4792 / 0.4769 ($\downarrow$ 0.48%)	0.4848 / 0.4822 ($\downarrow$ 0.54%)	0.4872 / 0.4864 ($\downarrow$ 0.16%)	0.4813 / 0.4783 ($\downarrow$ 0.62%)	0.4744 / 0.4672 ($\downarrow$ 1.52%)	0.4846 / 0.4786 ($\downarrow$ 1.24%)

Table 1: Comprehensive performance comparison. Rows denote datasets; columns denote models. Format: “Original / +TSDR (%Improv.)”.

ficacy: we observe the most significant relative lift in AKT, while DisKT achieves the best overall performance, likely attributed to its disentangled modeling and causal subtraction mechanism [Zhou *et al.*, 2025]. Furthermore, the positive results at $\gamma = 0.0$ indicate that our framework is “safe”, avoiding performance degradation even when selection bias is minimal, likely due to the regularization effect of the counterfactual imputation (see Appendix B for numerical results).

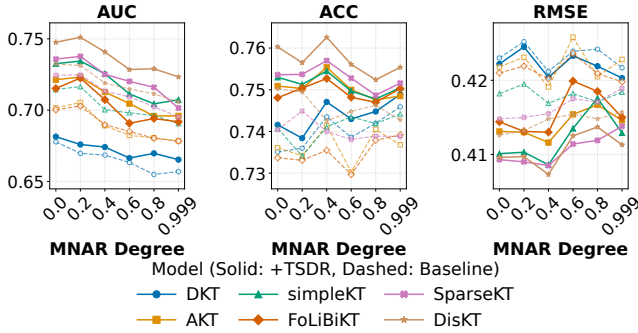
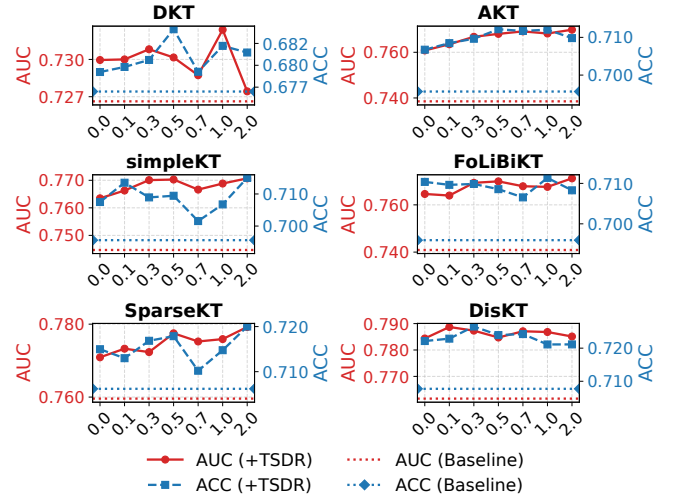


Figure 3: Performance of different MNAR levels on synthetic data.

Sensitivity Analysis (RQ3)

We focus on the prob dataset, where our model achieves superior performance, to highlight the experimental results. Figure 4 shows the impact of the hyperparameter λ on TSDR. We observe that incorporating the temporal smoothness constraint yields superior performance compared to the vanilla DR ($\lambda = 0$) within an appropriate range, while outperforming the original baselines. The method exhibits strong robustness, achieving a stable “sweet spot” generally in $\lambda \in [0.3, 1.0]$. Despite minor fluctuations in the case, e.g., DKT at $\lambda = 0.7$, the overall performance remains high, indicating that TSDR


 Figure 4: The AUC and ACC of TSDR on varying number of λ .

is effective without requiring meticulous tuning.

7 Conclusion

This paper introduces TSDR, a framework that addresses selection bias in knowledge tracing by integrating a doubly robust estimator with a theoretically-guided temporal smoothness constraint. Experiments on multiple benchmarks demonstrate its superior accuracy and stability over state-of-the-art methods. We acknowledge that the joint optimization of auxiliary modules introduces training overhead; however, this cost is confined to the offline phase and does not compromise online inference efficiency. Overall, TSDR offers a principled paradigm for learning from observational systems, paving the way for more reliable intelligent tutoring systems.

Acknowledgments

This research was supported in part by National Science and Technology Major Project (2021ZD0111501), Natural Science Foundation of China (U24A20233, 62206064, 62206061, 62476163, 62406078, 62406080), the Guangdong Basic and Applied Basic Research Foundation (2025A1515010172, 2023B1515120020), the Guangzhou Basic and Applied Basic Research Foundation (2024A04J4384), and CCF-DiDi GAIA Collaborative Research Funds (CCF-DiDi GAIA 202521).

Contribution Statement

Peilin Zhan, Wei Chen, and Weilin Chen contributed equally to this work and should be regarded as co-first authors. They jointly contributed to the conceptualization, methodology design, implementation, experimental analysis, and manuscript writing. All authors discussed the results and approved the final manuscript.

References

- [Anderson *et al.*, 1985] John R Anderson, C Franklin Boyle, and Brian J Reiser. Intelligent tutoring systems. *Science*, 228(4698):456–462, 1985.
- [Azuma, 1967] Kazuoki Azuma. Weighted sums of certain dependent random variables. *Tohoku Mathematical Journal, Second Series*, 19(3):357–367, 1967.
- [Baker *et al.*, 2008a] Ryan SJ d Baker, Albert T Corbett, and Vincent Aleven. More accurate student modeling through contextual estimation of slip and guess probabilities in bayesian knowledge tracing. In *International conference on intelligent tutoring systems*, pages 406–415. Springer, 2008.
- [Baker *et al.*, 2008b] Ryan SJd Baker, Albert T Corbett, and Vincent Aleven. Improving contextual models of guessing and slipping with a truncated training set. *Educational Data Mining*, page 67, 2008.
- [Baker, 2001] Frank B Baker. *The basics of item response theory*. ERIC, 2001.
- [Barrett *et al.*, 2024] Jake Barrett, Alasdair Day, and Kobi Gal. Improving model fairness with time-augmented bayesian knowledge tracing. In *Proceedings of the 14th Learning Analytics and Knowledge Conference*, pages 46–54, 2024.
- [Cai *et al.*, 2025] Ruichu Cai, Kaitao Zheng, Junxian Huang, Zijian Li, Zhengming Chen, Boyan Xu, and Zhifeng Hao. Causal view of time series imputation: Some identification results on missing mechanism. *International Joint Conferences on Artificial Intelligence*, 2025.
- [Carless, 2019] David Carless. Feedback loops and the longer-term: towards feedback spirals. *Assessment & Evaluation in Higher Education*, 44(5):705–714, 2019.
- [Chen *et al.*, 2024] Weilin Chen, Ruichu Cai, Zeqin Yang, Jie Qiao, Yuguang Yan, Zijian Li, and Zhifeng Hao. Doubly robust causal effect estimation under networked interference via targeted learning. In *International Conference on Machine Learning*, pages 6457–6485. PMLR, 2024.
- [Choi *et al.*, 2020] Youngduck Choi, Youngnam Lee, Dongmin Shin, Junghyun Cho, Seoyon Park, Seewoo Lee, Ji-neon Baek, Chan Bae, Byungsoo Kim, and Jaewe Heo. Ednet: A large-scale hierarchical dataset in education. In Ig Ibert Bittencourt, Mutlu Cukurova, Kasia Muldner, Rose Luckin, and Eva Millán, editors, *Artificial Intelligence in Education*, pages 69–73, Cham, 2020. Springer International Publishing.
- [Corbett and Anderson, 1994] Albert T Corbett and John R Anderson. Knowledge tracing: Modeling the acquisition of procedural knowledge. *User modeling and user-adapted interaction*, 4(4):253–278, 1994.
- [Cui *et al.*, 2023] Chaoran Cui, Hebo Ma, Chen Zhang, Chunyun Zhang, Yumo Yao, Meng Chen, and Yuling Ma. Do we fully understand students’ knowledge states? identifying and mitigating answer bias in knowledge tracing. *arXiv preprint arXiv:2308.07779*, 2023.
- [Cui *et al.*, 2025] Jiajun Cui, Hong Qian, Chanjin Zheng, Lu Wang, Mo Yu, and Wei Zhang. Rebalancing discriminative responses for knowledge tracing. *ACM Transactions on Information Systems*, 43(3):1–25, 2025.
- [de Haan *et al.*, 2019] Pim de Haan, Dinesh Jayaraman, and Sergey Levine. Causal confusion in imitation learning. *NeurIPS*, 2019.
- [Dudík *et al.*, 2011] Miroslav Dudík, John Langford, and Lihong Li. Doubly robust policy evaluation and learning. In Lise Getoor and Tobias Scheffer, editors, *Proceedings of the 28th International Conference on Machine Learning, ICML 2011, Bellevue, Washington, USA, June 28 - July 2, 2011*, pages 1097–1104. Omnipress, 2011.
- [Ghosh *et al.*, 2020] Aritra Ghosh, Neil Heffernan, and Andrew S Lan. Context-aware attentive knowledge tracing. In *Proceedings of the 26th ACM SIGKDD international conference on knowledge discovery & data mining*, pages 2330–2339, 2020.
- [Gorgun and Bulut, 2022] Guher Gorgun and Okan Bulut. Identifying aberrant responses in intelligent tutoring systems: an application of anomaly detection methods. *Psychological Test and Assessment Modeling*, 64(4):359–384, 2022.
- [Guo *et al.*, 2021] Xiaopeng Guo, Zhijie Huang, Jie Gao, Mingyu Shang, Maojing Shu, and Jun Sun. Enhancing knowledge tracing via adversarial training. In *Proceedings of the 29th ACM international conference on multimedia*, pages 367–375, 2021.
- [Huang *et al.*, 2023] Shuyan Huang, Zitao Liu, Xiangyu Zhao, Weiqi Luo, and Jian Weng. Towards robust knowledge tracing models via k-sparse attention. In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*, page 2441–2445, 2023.
- [Huang *et al.*, 2024] Changqin Huang, Hangjie Wei, Qionghao Huang, Fan Jiang, Zhongmei Han, and Xiaodi Huang. Learning consistent representations with temporal and causal enhancement for knowledge tracing. *Expert Systems with Applications*, 245:123128, 2024.

- [Im *et al.*, 2023] Yoonjin Im, Eunseong Choi, Heejin Kook, and Jongwuk Lee. Forgetting-aware linear bias for attentive knowledge tracing. In *Proceedings of the 32nd ACM international conference on information and knowledge management*, pages 3958–3962, 2023.
- [Li *et al.*, 2024] Xueyi Li, Youheng Bai, Teng Guo, Zitao Liu, Yaying Huang, Xiangyu Zhao, Feng Xia, Weiqi Luo, and Jian Weng. Enhancing length generalization for attention based knowledge tracing models with linear biases. In *Proceedings of the thirty-third international joint conference on artificial intelligence (IJCAI-24)*, pages 5918–5926, 2024.
- [Liu *et al.*, 2023] Zitao Liu, Qiongqiong Liu, Jiahao Chen, Shuyan Huang, and Weiqi Luo. simplekt: A simple but tough-to-beat baseline for knowledge tracing. In *The Eleventh International Conference on Learning Representations, ICLR 2023, Kigali, Rwanda, May 1-5, 2023*. OpenReview.net, 2023.
- [Liu *et al.*, 2025] Jie Liu, Zhifeng Li, Jiawang Yang, Hao He, and Fang Cui. The neurocognitive mechanism underlying math avoidance among math anxious people. *npj Science of Learning*, 10(1):50, 2025.
- [Lv *et al.*, 2025] Xiangwei Lv, Guifeng Wang, Jingyuan Chen, Hejian Su, Zhiang Dong, Yumeng Zhu, Beishui Liao, and Fei Wu. Debaised cognition representation learning for knowledge tracing. *ACM Transactions on Information Systems*, 2025.
- [Ma *et al.*, 2024] Haiping Ma, Yong Yang, Chuan Qin, Xiaoshan Yu, Shangshang Yang, Xingyi Zhang, and Hengshu Zhu. Hd-kt: Advancing robust knowledge tracing via anomalous learning interaction detection. In *Proceedings of the ACM Web Conference 2024*, pages 4479–4488, 2024.
- [Mao *et al.*, 2025] Shun Mao, Jieyu Zhan, Yuanfei Deng, Yixiu Qin, and Yuncheng Jiang. Improving exercise-level knowledge tracing via knowledge concept-based memory network. *Expert Systems with Applications*, page 127825, 2025.
- [Nakagawa *et al.*, 2019] Hiromi Nakagawa, Yusuke Iwasawa, and Yutaka Matsuo. Graph-based knowledge tracing: modeling student proficiency using graph neural network. In *IEEE/WIC/aCM international conference on web intelligence*, pages 156–163, 2019.
- [Newell and Rosenbloom, 2013] Allen Newell and Paul S Rosenbloom. Mechanisms of skill acquisition and the law of practice. In *Cognitive skills and their acquisition*, pages 1–55. Psychology Press, 2013.
- [Pandey and Karypis, 2019] Shalini Pandey and George Karypis. A self-attentive model for knowledge tracing. *International Educational Data Mining Society*, 2019.
- [Pearl *et al.*, 2016] J. Pearl, M. Glymour, and N.P. Jewell. *Causal Inference in Statistics: A Primer*. Wiley, 2016.
- [Piech *et al.*, 2015] Chris Piech, Jonathan Bassen, Jonathan Huang, Surya Ganguli, Mehran Sahami, Leonidas J Guibas, and Jascha Sohl-Dickstein. Deep knowledge tracing. *Advances in neural information processing systems*, 28, 2015.
- [Schnabel *et al.*, 2016] Tobias Schnabel, Adith Swaminathan, Ashudeep Singh, Navin Chandak, and Thorsten Joachims. Recommendations as treatments: Debiasing learning and evaluation. In *international conference on machine learning*, pages 1670–1679. PMLR, 2016.
- [Shen *et al.*, 2021] Shuanghong Shen, Qi Liu, Enhong Chen, Zhenya Huang, Wei Huang, Yu Yin, Yu Su, and Shijin Wang. Learning process-consistent knowledge tracing. In *Proceedings of the 27th ACM SIGKDD conference on knowledge discovery & data mining*, pages 1452–1460, 2021.
- [Steck, 2013] Harald Steck. Evaluation of recommendations: rating-prediction and ranking. In *Proceedings of the 7th ACM conference on Recommender systems*, pages 213–220, 2013.
- [Thomas and Brunskill, 2016] Philip Thomas and Emma Brunskill. Data-efficient off-policy policy evaluation for reinforcement learning. In *International conference on machine learning*, pages 2139–2148. PMLR, 2016.
- [Vaswani *et al.*, 2017] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017.
- [Villegas-Ch *et al.*, 2025] William Villegas-Ch, Diego Buenano-Fernandez, Alexandra Maldonado Navarro, and Aracely Mera-Navarrete. Adaptive intelligent tutoring systems for stem education: analysis of the learning impact and effectiveness of personalized feedback. *Smart Learning Environments*, 12(1):1–31, 2025.
- [Vygotsky *et al.*, 1978] Lev S Vygotsky, Michael Cole, Vera John-Steiner, S Scribner, and Ellen Souberman. The development of higher psychological processes, 1978.
- [Wang *et al.*, 2019] Xiaojie Wang, Rui Zhang, Yu Sun, and Jianzhong Qi. Doubly robust joint learning for recommendation on data missing not at random. In *International Conference on Machine Learning*, pages 6638–6647. PMLR, 2019.
- [Wise, 2017] Steven L Wise. Rapid-guessing behavior: Its identification, interpretation, and implications. *Educational Measurement: Issues and Practice*, 36(4):52–61, 2017.
- [Zambrano *et al.*, 2024] Andres Felipe Zambrano, Jiayi Zhang, and Ryan S Baker. Investigating algorithmic bias on bayesian knowledge tracing and carelessness detectors. In *Proceedings of the 14th Learning Analytics and Knowledge Conference*, pages 349–359, 2024.
- [Zhou *et al.*, 2025] Yiyun Zhou, Zheqi Lv, Shengyu Zhang, and Jingyuan Chen. Disentangled knowledge tracing for alleviating cognitive bias. In *Proceedings of the ACM on Web Conference 2025*, pages 2633–2645, 2025.