

**DEEP KNOWLEDGE TRACING WITH ATTENTION MECHANISM FOR
PERSONALIZED LEARNING PATH RECOMMENDATION IN INTELLIGENT
TUTORING SYSTEMS****Komronbek H. Obloev**

Asia International University

ABSTRACT

The advancement of intelligent tutoring systems (ITS) requires accurate modeling of student knowledge states over time. Traditional Bayesian Knowledge Tracing (BKT) and rule-based methods, while interpretable, suffer from limited expressive power and an inability to capture long-range dependencies between learning interactions. This study proposes an Attention-based Deep Knowledge Tracing (A-DKT) model that integrates a self-attention mechanism with a Long Short-Term Memory (LSTM) network to predict student mastery and recommend personalized learning paths in higher education environments.

The proposed architecture processes sequential learning interactions as input vectors enriched with skill embeddings, response correctness, and temporal features. The attention layer assigns dynamic weights to past interactions, enabling the model to identify which prior exercises most influence the current knowledge state. The output layer produces per-skill mastery probabilities that drive a path-recommendation module operating on top of the predicted knowledge state.

Experimental evaluation on a simulated dataset of 1,000 student interaction sequences shows that the proposed A-DKT model achieves an AUC of 0.88 and an RMSE of 0.29, outperforming standard DKT (AUC 0.81) and BKT (AUC 0.73). The model complements prior work on reinforcement-learning-based educational trajectory optimization by providing a more accurate state representation that the decision policy can act upon. Results indicate that attention-augmented knowledge tracing offers a practical, interpretable foundation for next-generation adaptive learning platforms.

Keywords: Deep Knowledge Tracing, Attention Mechanism, Intelligent Tutoring Systems, Personalized Learning, LSTM, Educational Data Mining, Adaptive Learning Paths.

1. INTRODUCTION

Personalized learning has become a defining objective of modern higher-education technology. The shift from one-size-fits-all instruction to individually adapted learning trajectories depends on a system's ability to maintain an accurate, up-to-date model of what each student knows. This task, formally known as knowledge tracing, sits at the core of every intelligent tutoring system (ITS): without a reliable estimate of the student's current mastery, neither content sequencing nor difficulty calibration nor remediation can be performed effectively.

Classical approaches to knowledge tracing rely on Bayesian Knowledge Tracing (BKT), a hidden Markov model in which a binary latent variable represents whether a student has learned a specific skill. While BKT is transparent and easy to deploy, it makes restrictive assumptions: skills are treated independently, learning is monotonic, and forgetting is typically ignored. These assumptions break down in realistic educational data, where skills overlap, students forget previously learned material, and engagement varies over time.

Deep Knowledge Tracing (DKT), introduced by Piech et al., reformulated the problem as a sequence-modeling task using recurrent neural networks. DKT demonstrated substantial gains over BKT but treats all prior interactions uniformly, which limits its ability to focus on the most informative parts of a learner's history. Recent advances in natural language processing have shown that attention mechanisms can address this limitation by assigning dynamic, context-dependent weights to elements of an input sequence.

In our previous work, we proposed an Adaptive Reinforcement Learning (ARL) framework that optimizes learning trajectories by treating instruction as a Markov Decision Process. That framework, however, assumed the availability of a high-quality state representation; the question of how to construct such a state from raw interaction logs was left open. The present study addresses precisely that gap. We propose an Attention-based Deep Knowledge Tracing (A-DKT) model that produces a fine-grained, per-skill mastery vector suitable for use as the state component of a downstream decision policy.

The objectives of this research are:

- To design an attention-augmented DKT architecture for sequential modeling of student knowledge states.
- To formalize a path-recommendation procedure that consumes the predicted mastery vector and selects the next instructional action.
- To evaluate the proposed model against BKT, standard DKT, and a hybrid DKT-ANFIS baseline on simulated learner data.
- To position the resulting model as a state-estimation module for adaptive reinforcement-learning-based tutoring systems.

2. RELATED WORK

Research on knowledge tracing can be grouped into three generations. The first generation, dominated by BKT and its extensions (BKT with item difficulty, individualized BKT, and BKT with forgetting), focused on interpretable parameters and small-scale deployments. These models remain popular in production tutoring systems because their behavior is easy to audit, but their predictive accuracy plateaus on large, noisy datasets.

The second generation, beginning with Deep Knowledge Tracing, applies recurrent neural networks (typically LSTM or GRU) to sequences of (skill, correctness) pairs. DKT and its variants — DKT+, Dynamic Key-Value Memory Networks (DKVMN), and Self-Attentive Knowledge Tracing (SAKT) — significantly outperform BKT on standard benchmarks such as ASSISTments and Statics2011. The trade-off is reduced interpretability: it becomes difficult to explain why the model predicts low mastery for a particular skill.

The third generation, emerging since 2022, combines deep models with structured knowledge: graph-based knowledge tracing (GKT), transformer-based knowledge tracing (AKT, SAINT), and hybrid neuro-symbolic models that integrate fuzzy reasoning or Bayesian priors. Our prior dissertation work on ANFIS-based student-performance modeling falls within this hybrid family. The model proposed in this paper sits in the second-to-third generation transition: it adds attention to DKT, which improves both accuracy and partial interpretability through visualizable attention weights, without introducing the engineering complexity of a full transformer stack.

3. METHODOLOGY

3.1 Problem Formulation

Let a student's interaction history be a sequence $x_1, x_2, \dots, x_t$, where each interaction $x_i = (s_i, c_i, \Delta t_i)$ consists of the skill identifier s_i , the correctness of the response $c_i \in \{0, 1\}$, and the time elapsed since the previous interaction Δt_i . The goal of knowledge tracing is to predict, for each future skill s , the probability that the student will respond correctly to the next exercise on that skill.

3.2 Input Representation

Each interaction is encoded as a concatenation of three embeddings:

- Skill embedding $e_s \in \mathbb{R}^d$, learned end-to-end during training.
- Response embedding $e_c \in \mathbb{R}^d$, encoding correctness and signaling that the exercise has been attempted.
- Temporal embedding $e_t \in \mathbb{R}^d$, a learned representation of the bucketed time interval Δt , used to capture forgetting effects.

The concatenated vector $x_i \in \mathbb{R}^{3d}$ is passed to the recurrent encoder.

3.3 Attention-Augmented Encoder

The encoder consists of a single-layer LSTM that produces hidden states $h_1, h_2, \dots, h_t$. A scaled dot-product attention layer is then applied to the sequence of hidden states with the current hidden state h_t as the query. The context vector c_t is computed as:

$$c_t = \sum_i \alpha_{ti} h_i, \quad \alpha_{ti} = \text{softmax}(h_t^T \cdot h_i / \sqrt{d})$$

The attention weights α_{ti} indicate how strongly each past interaction contributes to the current mastery estimate. This is the key mechanism by which the model can place more weight on, for example, a difficult exercise attempted ten steps ago than on a trivial recent one.

3.4 Output Layer

The concatenation $[h_t; c_t]$ is passed through a fully connected layer with sigmoid activation to produce a mastery vector $\hat{m}_t \in [0, 1]^K$, where K is the total number of skills. The component $\hat{m}_{t,k}$ is interpreted as the probability that the student would answer a next question on skill k correctly.

3.5 Training Objective

The model is trained to minimize the binary cross-entropy loss between the predicted probability for the actually answered skill and the observed correctness. Formally:

$$L = - \sum_t [c_t \log \hat{m}_{t,s_t} + (1 - c_t) \log (1 - \hat{m}_{t,s_t})]$$

The model is optimized with the Adam algorithm with an initial learning rate of 10^{-3} , gradient clipping at norm 5, and early stopping on a held-out validation set.

3.6 Path Recommendation Procedure

Given the predicted mastery vector $\hat{m}_t$, the recommender selects the next skill s^* using a simple but effective heuristic that balances mastery and prerequisite readiness:

- Filter out skills already mastered, i.e. $\hat{m}_{t,k} \geq \tau_{\text{high}}$ (default $\tau_{\text{high}} = 0.85$).
- Among the remaining skills, restrict to those whose prerequisites are sufficiently mastered, i.e. for all $p \in \text{prereq}(k)$, $\hat{m}_{t,p} \geq \tau_{\text{pre}}$ (default $\tau_{\text{pre}} = 0.70$).

- Within the eligible set, select the skill with mastery closest to a target band $[0.35, 0.65]$ — the zone of proximal development — to maximize learning gain per exercise.

This procedure produces a state-aware, prerequisite-respecting recommendation that can be passed to a higher-level reinforcement-learning policy as the action prior, or used directly as the final recommendation.

4. EXPERIMENTAL RESULTS

To evaluate the proposed model, we generated a simulated dataset of 1,000 student interaction sequences spanning 50 skills, with sequence lengths ranging from 40 to 200 interactions. The simulator implements a skill-graph with prerequisite relations, a per-skill learning rate, and a forgetting rate, producing data with characteristics similar to real ASSISTments-like logs while allowing controlled comparison. Eighty percent of the sequences were used for training, ten percent for validation, and ten percent for testing.

Four models were compared: classical BKT, standard DKT (single-layer LSTM with hidden size 64), a hybrid DKT-ANFIS model that post-processes DKT outputs through a fuzzy inference system, and the proposed A-DKT (LSTM with attention, hidden size 64, embedding dimension 32). The primary metric is the area under the ROC curve (AUC) for next-step correctness prediction; the secondary metric is RMSE between predicted and ground-truth mastery.

Table 1. Comparative evaluation of knowledge tracing models on the simulated dataset.

Model	Accuracy (AUC)	Prediction RMSE	Path Recommendation Quality
Bayesian Knowledge Tracing (BKT)	0.73	0.42	Low
Standard DKT (LSTM)	0.81	0.36	Medium
DKT + ANFIS (hybrid)	0.84	0.33	Medium-High
Proposed Attention-DKT	0.88	0.29	High

The Attention-DKT model achieves a 7-point AUC improvement over standard DKT and a 15-point improvement over BKT. Qualitative inspection of the attention weights confirms that the model learns to focus on transitions where the student first answered correctly after a prior failure on the same skill — a strong signal of learning. Path recommendations produced by feeding A-DKT outputs into the procedure described in Section 3.6 align with prerequisite ordering in 94% of cases, compared with 71% for the BKT baseline.

5. DISCUSSION

The experimental results support three main observations. First, the attention mechanism contributes a measurable accuracy gain on top of standard DKT, consistent with findings reported for SAKT and AKT on real-world benchmarks. The gain is largest in the early portion

of the sequence, where the model has fewer recent observations to draw on and benefits most from selectively attending to informative past interactions.

Second, the proposed model integrates naturally with our previously published ARL framework. The mastery vector $\hat{m}_t$ provides a continuous, K -dimensional state representation that replaces the coarse three-level (Low, Medium, High) discretization used in the prior MDP formulation. With a richer state, the Q-learning policy can express finer distinctions between actions and is expected to converge to a better policy. A direct comparison of joint A-DKT + ARL versus the original ARL is the subject of ongoing work.

Third, partial interpretability is restored. While the LSTM hidden state remains opaque, the attention weights α_{ti} form a one-dimensional explanation that can be surfaced to instructors: for any current recommendation, the system can report which past three or four interactions most influenced the decision. This addresses a recurring concern about deep tutoring systems in higher-education deployment, where instructors are reluctant to act on recommendations they cannot audit.

Several limitations should be noted. The simulator, while parameterized to match published statistics, cannot fully replicate the noise and metacognitive behavior of real students; validation on an authentic dataset is the immediate next step. The current model treats skills as discrete identifiers; richer skill representations grounded in domain ontologies would likely improve cold-start performance on previously unseen skills. Finally, the path recommender uses fixed thresholds; learning these thresholds — for instance, through reinforcement learning, closing the loop with our earlier framework — is a natural extension.

6. CONCLUSION

This study presented an Attention-based Deep Knowledge Tracing model designed to serve as a state-estimation module within adaptive intelligent tutoring systems. By augmenting an LSTM encoder with scaled dot-product attention and coupling its output to a prerequisite-aware path recommender, the proposed system achieves substantially higher predictive accuracy than BKT, standard DKT, and a DKT-ANFIS hybrid, while preserving a degree of interpretability through visualizable attention weights.

The model complements our prior reinforcement-learning framework for personalized educational trajectories: where the earlier work optimized the decision policy, the present work supplies the high-resolution state on which the policy depends. Together, the two components form an end-to-end pipeline for adaptive higher-education platforms — from raw interaction logs, through fine-grained knowledge estimation, to next-action recommendation.

Future work will focus on three directions: (1) validation on authentic university-scale interaction datasets; (2) joint end-to-end training of A-DKT with the ARL policy, replacing the current modular pipeline; and (3) extension to multimodal inputs, incorporating engagement signals such as response time and revisit patterns to further enrich the knowledge state.

REFERENCES

1. Piech, C., Bassen, J., Huang, J., Ganguli, S., Sahami, M., Guibas, L., & Sohl-Dickstein, J. (2015). Deep Knowledge Tracing. *Advances in Neural Information Processing Systems*, 28.
2. Pandey, S., & Karypis, G. (2019). A Self-Attentive Model for Knowledge Tracing. *Proceedings of the 12th International Conference on Educational Data Mining (EDM)*.
3. Ghosh, A., Heffernan, N., & Lan, A. S. (2020). Context-Aware Attentive Knowledge Tracing. *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*.
4. Corbett, A. T., & Anderson, J. R. (1995). Knowledge Tracing: Modeling the Acquisition of Procedural Knowledge. *User Modeling and User-Adapted Interaction*, 4(4), 253–278.
5. Zhang, J., Shi, X., King, I., & Yeung, D.-Y. (2017). Dynamic Key-Value Memory Networks for Knowledge Tracing. *Proceedings of the 26th International Conference on World Wide Web*.
6. Vaswani, A., Shazeer, N., Parmar, N., et al. (2017). Attention Is All You Need. *Advances in Neural Information Processing Systems*, 30.
7. Obloev, K. H. (2026). Adaptive Reinforcement Learning Framework for Personalized Educational Trajectories in AI-Driven Higher Education Systems.
8. Holmes, W., Bialik, M., & Fadel, C. (2019). *Artificial Intelligence in Education*. Center for Curriculum Redesign.
9. Sutton, R. S., & Barto, A. G. (2018). *Reinforcement Learning: An Introduction* (2nd ed.). MIT Press.
10. Shin, D., Shim, Y., Yu, H., Lee, S., Kim, B., & Choi, Y. (2021). SAINT+: Integrating Temporal Features for EdNet Correctness Prediction. *Proceedings of the 11th International Conference on Learning Analytics and Knowledge*.