

Optimizing Feedback Loops in Adaptive Learning Systems: A Computational Model-Based Integration of Diagnostic Analytics and Cognitive Load Principles

Mr. Amirreza Mehrabi, Purdue Engineering Education

Dr. Amirreza Mehrabi is an accomplished scholar at Purdue University specializing in machine learning and AI-aided modeling and Engineering Education. As an interdisciplinary researcher, he designs and deploys intelligent platforms that tackle complex decision-making challenges, integrating advanced techniques such as transformer models, reinforcement learning, and statistical inference. His work spans across adaptive testing, cognitive fatigue detection by ML, agentic AI agents tutors, and multimodal analytics, leading to innovative solutions with measurable real-world impact. Amirreza's expertise is reflected in his custom architectures, model optimization research, and psychometrics—combining educational measurement science with deep learning. Amirreza's collaborative spirit, technical mastery, and commitment to transparent, bias-aware modeling foster a vibrant culture of innovation.

Dr. Jason W Morphew, Purdue University – West Lafayette

Dr. Jason Morphew is an assistant professor at Purdue University in Engineering Education and serves as the director of undergraduate curriculum and advanced learning technologies for SCALE, and is affiliated with the INSPIRE research institute for Pre-College Engineering and the Center for Advancing the Teaching and Learning of STEM. Dr. Morphew's research focuses on the application of principles of learning derived from cognitive science and the learning sciences to the design and evaluation of learning environments and technologies that enhance learning, interest, and engagement in STEM. Morphew earned his Ph.D. from the University of Illinois Urbana-Champaign in Educational Psychology from the Cognitive Science of Teaching and Learning Division where he examined self-regulated learning of engineering students enrolled in introductory physics courses.

Breejha Sene Quezada, Purdue University

Prof. N. Sanjay Rebello

Dr. Sanjay Rebello is Professor of Physics and Astronomy and Professor of Curriculum and Instruction at Purdue University. His research focuses on the teaching and learning of physics. He is particularly interested in issues pertaining to transfer of le

Optimizing Feedback Loops in Adaptive Learning Systems: A Computational Model-Based Integration of Diagnostic Analytics and Cognitive Load Principles

Amirreza Mehrabi^{1,2} Jason W. Morphew¹

Breejha Quezada³ N. Sanjay Rebello⁴

¹School of Engineering Education, Purdue University

²Dept. of Computer and Electronic Engineering, Purdue University

³School of Business and Information Technology, Purdue Global

⁴Department of Physics and Astronomy, Purdue University

amehrabi@purdue.edu

Abstract

While adaptive learning systems in engineering education show acceptable performance in diagnostic reports, transforming these reports into actionable decision-making evidence remains underexplored. This study focuses explicitly on the feedback recommendation control problem, in which students receive diagnostic reports and require targeted review materials drawn from existing resources, where refining existing knowledge constructs is as important as initial instruction. This study designs a closed-loop architecture that integrates Item Response Theory (IRT), Cognitive Diagnostic Modeling (CDM), and Large Language Models (LLMs) to support feedback-mediated learning under cognitive constraints. The IRT model assumes learning as a latent trait, and the CDM assumes learning as the binary mastery of discrete learning objects, both by estimating item characteristics. These characteristics are difficulty, guessing, and slipping on the items (questions). LLMs contextualize diagnostic evidence by aligning identified skill gaps with course-specific learning resources. Using greedy and gradient-based optimization, we model the integration of diagnostic reports with available learning resources to select materials that address identified skill gaps while controlling cognitive load and redundancy. The framework was evaluated using 1,000 simulated learners and 720 engineering students. The framework was further validated through expert review, showing 82% agreement between model and expert-selected resources for engineering students. Results demonstrate that in simulation, across both GREEDY_LLM and GRADIENTDESCENT_LLM, 100% coverage was achieved, while in the real classroom deployment, GREEDY_LLM maintained 100% coverage at an average of 3.9 resources per student, and GRADIENTDESCENT_LLM achieved stronger semantic alignment and higher overall utility at an average of 7.5 resources per student with 97.5% coverage. Therefore,

this article indicates the feasibility of adaptive feedback operationalized as a controlled and bounded instruction.

Introduction

In engineering education, post-assessment feedback targets partial or inconsistent knowledge, not an absence of knowledge (Lipnevich & Panadero, 2021; Shute, 2008). Despite this, classroom and platform implementations still treat mistakes at the item level. After an incorrect response, teachers and platforms typically provide the correct solution. However, they often do not investigate the cause of the mistake or examine the underlying construct that led to the error. This approach risks reducing misunderstandings to isolated item-level errors, leaving underlying misconceptions or missing constructs unaddressed. (Desmarais & Baker, 2012; Lipnevich & Panadero, 2021). Because of this, errors from the same learning objectives, skills, or underlying constructs tend to recur (Desmarais & Baker, 2012; Lipnevich & Panadero, 2021; Mehrabi et al., 2024). Remediation requires mapping diagnostic evidence to bounded instructional actions—such as recommending learning resources that target the missing or weak skill(s), which can be learning objectives or any other type of underlying construct—while respecting working-memory limits and cognitive fatigue (Ouwehand et al., 2025; Paas et al., 2003).

Item Response Theory (IRT) estimates a continuous latent proficiency that is invariant to the administered items, separating learner ability from item characteristics such as difficulty (Fan et al., 2021). Cognitive diagnostic models (CDMs) complement this by inferring discrete learning objects as groups of invariant learning objectives that depend on the administered items with their specific item characteristics (Mehrabi et al., 2024, 2025). However, diagnosis alone does not determine how to *structure* feedback. In particular, many adaptive systems follow CDM and IRT concepts, which make them able to identify that an answer is wrong, and as remediation prefer local approaches such as providing written feedback on that item in a different format or asking similar questions (Desmarais & Baker, 2012; Lipnevich & Panadero, 2021). As instructional resources vary in scope and prerequisite alignment, broad remediation can increase cognitive load (Karaman, 2021; Paas et al., 2003; Shute, 2008). At the same time, very specific resources that do not address the underlying understanding can produce a similar outcome, as students may only be able to answer one question by memorizing without repairing the underlying construct (Ouwehand et al., 2025; Paas et al., 2003).

This study addresses two research questions: **RQ1:** How can post-assessment diagnostic evidence be systematically converted into targeted instructional feedback in engineering courses? **RQ2:** To what extent can the reviewing resources be bounded to cover all diagnosed skill gaps while minimizing cognitive load for the learner?

We frame adaptive learning after assessment as a closed-loop control problem: (i) providing the diagnostic reports of the students from their assessment, (ii) understanding the item- and content-level information via all student responses to the questions and learning resource information, and (iii) relying on an optimization-based selection that addresses the problem of selecting the minimum amount of resources that target the missing learning objective or weak underlying construct (Lipnevich & Panadero, 2021; Mehrabi et al., 2025).

Literature Review

Feedback in class: refinement, not reteaching

Misconceptions typically hold partially correct, fragmented, or inconsistent knowledge structures (Das, 2023; Litzinger et al., 2011). This common phenomenon can be observed in engineering education as concepts become more demanding and require more cognitive load. In such a situation, feedback must function as conceptual refinement, repairing and stabilizing knowledge structures so they support transfer and subsequent learning (Das, 2023; Litzinger et al., 2011). Longitudinal research links effective feedback to a shift from surface strategies that try to locally solve the problem without addressing the underlying issue, like hints on one question, toward deeper learning behaviors such as conceptual integration (Litzinger et al., 2011). At the same time, When addressing misconceptions through feedback, the goal is not simply to provide additional information but to organize existing knowledge in ways that reduce cognitive load during extended review (Joshi & Davis, 2024). Cognitive load theory makes this constraint explicit: feedback or any learning resource should be bounded in scope and targeted to reduce extraneous load (Paas et al., 2003). Thus, classroom feedback is not merely an item-level intervention; it is an instructional intervention that tackles the underlying misconception or knowledge gap via selective, pre-organized learning resources that exactly address the underlying construct and do not add cognitive load to students (Paas et al., 2003). This challenge is further amplified in contexts where access to instructor support is limited, such as in large or first-year engineering courses.

The diagnostic–action gap

Statistically principled models for representing learner state, including IRT, Bayesian networks, and related latent-variable approaches, have been well established in the literature (Bichi & Talib, 2018; Desmarais & Baker, 2012). These models treat performance as a sequence of items answered by a group of examinees and attempt to estimate proficiency, mastery, and uncertainty based on their assumptions about the underlying construct (Desmarais & Baker, 2012). While these models primarily address the measurement question of “What is the learner state?” they generally do not provide guidance on the instructional action to take (Desmarais & Baker, 2012). Indeed, diagnostic reports do not automatically provide classroom-useful feedback that targets the causal missing attribute (Desmarais & Baker, 2012).

The existing feedback generators follow the recommendation system organization, which provides related resources without guaranteeing prerequisite coverage (Pinos Ullauri, 2024). Indeed, they do not follow the underlying construct assumptions of models like IRT and CDM (Pinos Ullauri, 2024). The gap that we aim to address is that classroom feedback requires a principled mapping from diagnostic evidence to actions that target missing skill gaps under cognitive load constraints (Desmarais & Baker, 2012; Paas et al., 2003).

Methodology

We designed a pipeline to separate what the learner knows from what resources are allocated. The learning profile is built from the responses of the students and the learning resource profile is built

from the information of the resource. The pipeline is shown in Fig. 1: (i) responses Y_{ij} feed an IRT and CDM (Deterministic Input, Noisy “And” Gate (DINA) model) that estimates performance as a continuous and class-based measure, which we call overall proficiency (θ_i) and mastery on learning objectives or skills (S_{ik}); (ii) instructional content is augmented by features from a frozen LLM encoder; and (iii) an optimization engine selects the best resource based on each learner’s proficiency profile to maximize skill coverage and minimize review time (Mehrabi et al., 2025; Paas et al., 2003).

We simulated $N = 1,000$ students responding to $I = 60$ items associated with $K = 5$ skills using the IRT 3PL model. The Q-matrix as the matrix that explains the item and skill/learning objectives is described as $Q \in \{0, 1\}^{I \times K}$ (Appendix; Table S1). To be more realistic, 60% of items target a single skill and 40% require more than one skill, up to three. Student mastery profiles were defined as $\alpha_n \in \{0, 1\}^K$ with $\alpha_{nk} \sim \text{Bernoulli}(0.6)$. Slipping and guessing parameters were simulated as $s_i \sim \text{Beta}(5, 15)$ with mean ≈ 0.25 and $g_i \sim \text{Beta}(7, 18)$ with mean ≈ 0.28 . Responses were generated by the DINA model, with $\eta_{ni} = \prod_{k=1}^K \alpha_{nk}^{Q_{ik}}$ indicating whether student n has mastered all skills required by item i . The probability of a correct response is

$$P(y_{ni} = 1 \mid \eta_{ni}) = g_i^{1-\eta_{ni}} (1 - s_i)^{\eta_{ni}},$$

and y_{ni} was sampled by comparing this probability to $u \sim \text{Uniform}(0, 1)$.

Simulated instructional repositories. We created a synthetic repository of $M \in \{5, 10, 15, 20\}$ resources assigned to $K = 5$ skills with a binary matrix $C \in \{0, 1\}^{M \times K}$. Resource durations L_r were generated by a truncated log-normal distribution bounded at 5–15 minutes,

$$L_r \sim \max \left(\min(\text{LogNormal}(\mu = \log(20), \sigma = 2), 15), 5 \right),$$

to mimic realistic learning resource durations (Paas et al., 2003). Resources were divided into 30% easy, 50% medium, and 20% hard, with 80% of resources covering only one skill and 20% cover two randomly chosen skills.

The real-world classroom assessment included $I = 19$ items representing $K = 4$ course-specific skills and learning objectives. The dataset consisted of 720 students and was used for model comparison in Tables 1–4. The optimizer uses the empirical C_{rk} obtained by semantically annotating each resource with difficulty and a skill-coverage vector over course skills and learning objectives including systems and trigonometric setup, torque, total angular momentum, and principle of angular momentum. The complete content–skill matrix is located in the Appendix.

Step 1: Dual Diagnosis

The student performance profile is diagnosed with two models. A 3PL IRT model estimates a continuous proficiency parameter θ_i and item parameters of difficulty, discrimination, and guess (a_j, b_j, c_j) (Desmarais & Baker, 2012). The DINA model generates a binary mastery profile S_{ik} over skills, in line with the course Q-matrix (Desmarais & Baker, 2012). The set of non-mastered skills,

$$\mathcal{G}_i = \{k : S_{ik} = 0\},$$

represents the particular skill gaps that feedback must address.

Step 2: Resource model and bounded selection (optimization)

Each resource r has an estimated time cost t_r , a redundancy penalty d_r for the overlap with other selected resources, and a binary indicator of coverage C_{rk} indicating whether resource r covers skill k . We use the diagnostic evidence to elect a minimal subset of resources that covers $\mathcal{G}_i$ while controlling cognitive load and bounds the cognitive load (Mehrabi et al., 2025). Let selection be $x_r \in \{0, 1\}$. We solve for learner i :

$$\min_{\mathbf{x}} \sum_r x_r (t_r + \lambda d_r) - \mu \sum_{k \in \mathcal{G}_i} \log \left(\epsilon + \sum_r x_r C_{rk} \right). \quad (1)$$

The small constant $\epsilon > 0$ is added inside the log to ensure numerical stability during optimization, particularly in intermediate iterations where coverage may be incomplete (Paas et al., 2003; van Merriënboer & Ayres, 2005). We penalize the selection of redundant content once a skill has been covered and enforce prerequisite constraints such that resources requiring unmastered prerequisite skills are not selected (Desmarais & Baker, 2012).

Step 3: Semantic alignment

This work uses semantic similarity as an auxiliary content feature to better align selected resources with identified learner skill gaps. This consideration covers scenarios where students missed questions that the errors arise from contextual similarities not fully captured by the predefined skill mapping to the learning content (Mehrabi et al., 2025). We leverage a pretrained LLM, specifically Mistral Small 3.1, as a semantic encoder to represent assessment items and instructional resources and calculate cosine similarities s_{ir} between item i and resource r . The Q-matrix aggregates similarities to the skill level,

$$S_{rk} = \frac{1}{|I_k|} \sum_{i \in I_k} s_{ir}, \quad (2)$$

which is used to modulate coverage via $\tilde{C}_{rk} = C_{rk} \cdot S_{rk}$.

Step 4: Solution procedure

We take two different approaches. A greedy approach for interpretability chooses the resource at each step that maximizes the uncovered skill coverage per cognitive load:

$$r^* = \arg \max_{r \in \mathcal{R}} \frac{\sum_{k \in \mathcal{G}_i} \tilde{C}_{rk}}{t_r + \lambda d_r}, \quad (3)$$

and stops once all skills in $\mathcal{G}_i$ are covered (Mehrabi et al., 2025). For more heterogeneous libraries, we relax $x_r \in [0, 1]$ and optimize the same objective with gradient-based methods using $\tilde{C}_{rk}$ and discretize the resulting set (Mehrabi et al., 2025).

We conducted an expert validation study where three domain experts independently reviewed a random subset of 10 selected engineering students. Experts selected videos that they considered appropriate for each student’s skill gaps, without being informed of the model’s selections.

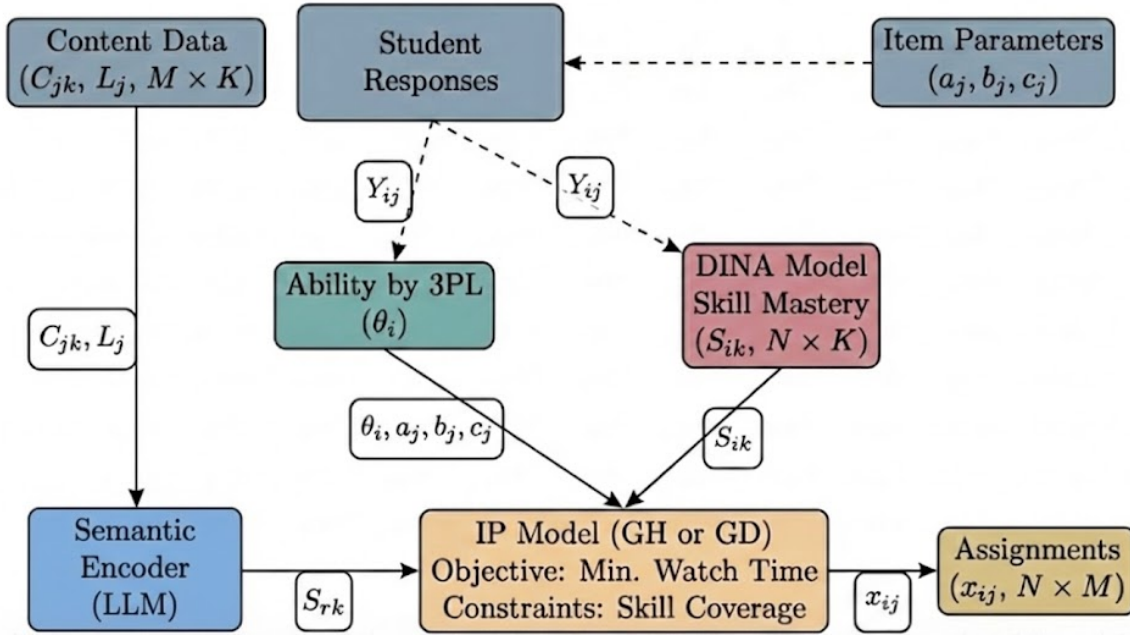


Figure 1: **Closed-loop adaptive feedback pipeline.** Student responses feed dual diagnosis (3PL IRT proficiency and DINA skill mastery). Content metadata (coverage, time, difficulty) and LLM-derived semantic alignment enter a constrained selection module that produces bounded assignments with coverage and prerequisite guarantees.

Agreement between expert-selected resources and model-selected resources was then measured. This agreement was computed as the proportion of overlapping selected resources between the model and expert selections for each student.

Results

The differences between the two optimization models reflect the mapping from skill gaps to the assigned resource set. We compare the models on the basis of whether diagnosed skill gaps are fully covered, whether remediation is bounded in cognitive load, and whether higher objective scores represent more efficient decisions or simply more content. Appendix Tables S1–S13 provide supporting materials for our optimization model performance.

Skill coverage

A model covering all skills in each student’s diagnosed skill gap set means the algorithm achieves full coverage of the diagnosed skill gaps. It should be noted that Table of S2, also indicate the item level information. Table 1 shows a clear separation between the two optimization models. GREEDY_LLM covers all required skills across all students with zero under-coverage. GRADIENTDESCENT_LLM covers all required skills in most cases but misses at least one skill for 18 students, which is 2.5% of cases.

However, the amount of missing skills under GRADIENTDESCENT_LLM is limited. Table 2

Table 1: Coverage composition by method

Method	Exact_Fit	Over_Coverage	Under_Coverage	No_Coverage	No_Requirement
GRADIENTDESCENT_LLM	381 (52.9%)	321 (44.6%)	18 (2.5%)	0 (0.0%)	0 (0.0%)
GREEDY_LLM	293 (40.7%)	427 (59.3%)	0 (0.0%)	0 (0.0%)	0 (0.0%)

Table 2: Severity of coverage violations.

Method	Mean_Missing	Max_Missing	Missing_eq_1_pct	Missing_ge_2_pct	Mean_Extra
GREEDY_LLM	0.000	0	0%	0%	2.194
GRADIENTDESCENT_LLM	0.025	1	2.5%	0%	3.690

indicates that all cases with missing skills miss exactly one required skill and no cases miss two or more skills. This is confirmed in Appendix Table S6, where the missing-skill measure is mostly zero, indicating that missing skills are a rare boundary behavior, not a systematic failure. The main difference in the behavior of the two models is redundancy: on average, GRADIENTDESCENT_LLM has more extra skills and the dispersion of extra coverage is much wider than GREEDY_LLM. This allows us to distinguish a model that always covers the required skills from a model that explores a broader neighborhood of solutions that includes both increased redundancy and occasional missing skills.

According to Table , skill coverage is uneven across skills. In particular, one skill is supported by only a single resource, while others are supported by multiple alternatives. This bottleneck helps explain the small number of under-coverage cases observed under GRADIENTDESCENT_LLM.

Boundedness: objective gains versus intervention burden

The central trade-off for classroom feedback is whether improvements in objective value and semantic alignment come from more efficient resource sets or from expanding remediation volume. Table 4 shows that GRADIENTDESCENT_LLM achieves higher mean utility and higher semantic alignment, including nearly double the total similarity mass per student. However, these gains come with substantially greater cognitive load as GRADIENTDESCENT_LLM assigns almost twice as many resources per student compared to GREEDY_LLM. Because of this, the marginal value per resource drops: utility per resource drops from 0.651 under GREEDY_LLM to 0.420 under GRADIENTDESCENT_LLM. This is the same pattern we see in cognitive load situations: objective gains are largely purchased by assigning more content rather than selecting a comparably sized but more efficient resource set (Paas et al., 2003). Appendix Table S5 shows that GRADIENTDESCENT_LLM maintains higher utility across the distribution but exhibits wide dispersion in assigned resource counts, whereas GREEDY_LLM concentrates interventions into smaller resource sets with higher marginal returns. In practical terms, GREEDY_LLM selects fewer resources while covering all required skills. GRADIENTDESCENT_LLM assigns more resources to improve alignment scores even when marginal gains per resource diminish, and occasionally misses a required skill.

Table 3: Instructional resources, skill coverage indicators, and difficulty levels.

Material ID	Skill 1	Skill 2	Skill 3	Skill 4	Difficulty
19_2	0	1	0	0	Basic
19_3	0	1	0	0	Medium
20_1	0	0	0	1	Medium
20_2	0	1	0	0	Basic
20_3	0	1	0	0	Hard
20_4	0	0	1	0	Medium
21_1	0	0	0	1	Basic
21_3	0	0	1	0	Medium
22_1	0	0	0	1	Medium
22_2	0	0	0	1	Medium
22_3	0	0	0	1	Hard
22_4	1	1	0	1	Hard
23_1	0	1	0	1	Medium
23_2	0	1	0	1	Medium
24_1	0	0	1	0	Medium

Robustness and mechanism: same targets, different implementations

GREEDY_LLM always covers the required skills and assigns fewer resources, whereas GRADIENTDESCENT_LLM consistently assigns larger resource sets while maintaining small but non-zero missing-skill rates. At the same time, GRADIENTDESCENT_LLM exhibits greater variation in assigned resources and alignment across quartiles, which suggests higher responsiveness to proficiency, while GREEDY_LLM maintains a more stable intervention volume.

Both models correctly identified skills to address; however, they vary in which resources they use to cover those skills. GRADIENTDESCENT_LLM distributes assignments across a wider set of resources as shown in Appendix Table S8, whereas GREEDY_LLM repeatedly selects a smaller subset of high-leverage resources. Skill-level aggregation is nearly identical across both models as shown in Appendix Table S9, which means both models agree on what skills to target and differ primarily in which resources they use. Paired within-student comparisons reinforce this: overlap in covered skills is high with a median Jaccard of 1.0 while overlap in selected resources is much lower with a median Jaccard of 0.20 as shown in Appendix Table S10. When GRADIENTDESCENT_LLM improves utility relative to GREEDY_LLM, it does so predominantly by assigning more resources in 87.8% of paired cases rather than finding a more efficient resource set of similar size as shown in Appendix Table S11. Missing required skills under GRADIENTDESCENT_LLM remains rare and always represents a worsening relative to GREEDY_LLM as shown in Appendix Table S12. Because of this, this highlights the importance of closing the diagnostic–action gap by ensuring full skill coverage under resource constraints. Across the evaluated cases, expert selections matched the model’s recommended resources in 83% of instances, with 83% agreement on assigned resource duration. This level of agreement

Table 4: **Core decision-quality metrics.** Student-level means (same diagnostic inputs for both models).

Metric	GRADIENTDESCENT_LLM	GREEDY_LLM
Students	720	720
Satisfactory_Rate	0.975	1.000
CoverageRate_mean	0.996	1.000
Utility_mean	2.988	2.443
AssignedResources_mean	7.454	3.901
UtilityPerResource_mean	0.420	0.651
LLMsim_mean	0.665	0.626
LLMsim_sum_mean	4.837	2.405
Missing_mean	0.025	0.000
Extra_mean	3.690	2.194

suggests that the model’s recommendations are largely aligned with expert judgment, providing qualitative support for the validity of the resource selection approach (Appendix A, Table S4).

Discussion

Learning in engineering domains is cumulative, with later concepts building on earlier ones. If students advance without mastering a prerequisite, the deficit accumulates over time and is more difficult to address later (Desmarais & Baker, 2012). Hence, we argue that full skill coverage should be a hard constraint—something that the system must guarantee, not a soft preference to be traded off with other desirable qualities during optimization. GRADIENTDESCENT_LLM achieves better scores on alignment metrics, but it does so not by finding a more targeted selection, but rather by assigning more resources. Students entering the remediation phase typically possess partial or incomplete knowledge rather than starting from scratch (Lipnevich & Panadero, 2021). Adding more resources that are less relevant increases cognitive load without proportional learning gain, undermining the purpose of the feedback. Semantic alignment is a useful criterion, but only as a tie-breaker within the space of solutions that already satisfy the hard coverage constraint (Desmarais & Baker, 2012; Lipnevich & Panadero, 2021). Optimizing for alignment without ensuring coverage and bounded cognitive load does not improve feedback; it only improves scores through means that do not lead to better learning outcomes.

To address RQ1, we diagnosed the skills a student lacked based on their exam performance and selected the remedial resources to assign based on that diagnosis. This separation allows us to formalize the feedback process and evaluate it rigorously through two LLM-based models, GREEDY_LLM and GRADIENTDESCENT_LLM, which take the diagnostic output as input and produce a concrete set of resource recommendations that address the identified skill gaps, based on prerequisite structure and semantic alignment with learning objectives (Desmarais & Baker, 2012; Lipnevich & Panadero, 2021). For RQ2, we demonstrate that resource bounding is both feasible and necessary, but that it must be enforced intentionally, not as a natural outcome of

optimization. GREEDY_LLM is able to cover all diagnosed skill gaps while minimizing the number of resources assigned. GRADIENTDESCENT_LLM demonstrates the effect of not enforcing the bound as a hard constraint, tending to assign more resources to improve alignment scores even when the extra content provides diminishing value. Since students in the remediation phase have partial knowledge, too broad a resource set can lead to cognitive overload with no proportional learning gain, confirming that full coverage and bounded cognitive load are compatible only if coverage is a hard constraint and resource count is explicitly constrained (Desmarais & Baker, 2012; Lipnevich & Panadero, 2021). Results highlight that coverage performance depends not only on the optimization strategy but also on the structure of the instructional repository. Ensuring multiple resources per skill is therefore critical for robust adaptive feedback systems.

Conclusion

The most effective post-assessment feedback targets partial knowledge through feasible and cognitively bounded instruction (Paas et al., 2003; Shute, 2008). Reformulating remediation as a constrained control problem reveals the diagnosis-to-action mapping as a crucial design decision. For engineering students, GREEDY_LLM satisfies diagnosed skill coverage requirements with fewer resources, while GRADIENTDESCENT_LLM improves objective scores through increased resource allocation, highlighting a trade-off between alignment and efficiency (Mehrabi et al., 2025). Closing the diagnostic-action gap in classroom feedback means enforcing full skill coverage and bounded cognitive load as hard constraints, then using semantic alignment to select among valid solutions. In settings where the library of resources is large or weakly organized, GRADIENTDESCENT_LLM can outperform by coordinating trade-offs globally across resources (Mehrabi et al., 2025). However, alignment scores can also be achieved at the cost of increased cognitive load due to the longer review time required (Paas et al., 2003; Shute, 2008).

Real-world application

This system shifts the focus of review from a generic review to a targeted review that requires the minimum time to identify and study the most relevant parts of the resources, rather than reviewing all content regardless of its relation to the diagnosed skill gaps. In an implementation setting, instructors first label their resources in a structured chain, specifying how each item relates to particular labels. These labels can be learning objectives or skills that link each question to the relevant resource. The GitHub repository (<https://github.com/amehrabi67/Video-content-selection-by-DINAandIRT-Optimized-by-GradientDec-GreedyHuristic>) requires three inputs: a student-item response matrix containing binary performance data, a Q-matrix mapping assessment items to specific latent skills, and a resource-skill matrix detailing resource coverage, duration, and difficulty levels. The content selection algorithm rapidly optimizes the selection process to minimize cognitive load by balancing total instructional duration with specified difficulty and a full diagnostic report of students. In general, for introductory engineering courses, where the resources are more controlled and focused, GREEDY_LLM is recommended; in upper-level courses, with richer, more diverse, and more heterogeneous resources, GRADIENTDESCENT_LLM provides better support for content selection.

Acknowledgments

This material is based upon work supported by the National Science Foundation under Awards No. 2322015 and No. 2142317. Any opinions, findings and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

References

- Bichi, A. A., & Talib, R. (2018). Item response theory: An introduction to latent trait models to test and item development. *International Journal of Evaluation and Research in Education*, 7(2), 142–151. <https://doi.org/10.11591/ijere.v7.i2.pp142-151>
- Das, D. K. (2023). Exploring the impact of feedback on student performance in undergraduate civil engineering. *European Journal of Engineering Education*, 48(6), 1148–1164.
- Desmarais, M. C., & Baker, R. S. J. d. (2012). A review of recent advances in learner and skill modeling in intelligent learning environments. *User Modeling and User-Adapted Interaction*, 22(1-2), 9–38.
- Fan, T., Yan, X., Song, J., & Guan, Z. (2021). Diagnosing english reading ability in chinese senior high schools. *Language Testing in Asia*, 11(2).
- Joshi, S. S., & Davis, K. A. (2024). Exploring shifts in engineering students' understanding of engineering work during their co-op experiences. *Proceedings of REES 2024*.
- Karaman, P. (2021). The effect of formative assessment practices on student learning: A meta-analysis study. *International Journal of Assessment Tools in Education*, 8(4), 801–817. <https://doi.org/10.21449/ijate.870300>
- Lipnevich, A. A., & Panadero, E. (2021). A review of feedback models and theories: Descriptions, definitions, and conclusions. *Frontiers in Education*, 6. <https://doi.org/10.3389/feduc.2021.720195>
- Litzinger, T., Lattuca, L. R., Hadgraft, R., & Newstetter, W. (2011). Engineering education and the development of expertise. *Journal of engineering education*, 100(1), 123–150.
- Mehrabi, A., Morphew, J. W., Araabi, B. N., Memarian, N., & Memarian, H. (2024). Ai-enhanced decision-making for course modality preferences in higher engineering education during the post-covid-19 era. *Information*, 15(10), 590.
- Mehrabi, A., Morphew, J. W., Quezada, B., & Rebello, N. S. (2025). Making evidence actionable in adaptive learning. *arXiv preprint arXiv:2511.14052*.
- Ouwehand, K., Lespiau, F., Tricot, A., & Paas, F. (2025). Cognitive load theory: Emerging trends and innovations. *Education Sciences*, 15(4), 458. <https://doi.org/10.3390/educsci15040458>
- Paas, F., Renkl, A., & Sweller, J. (2003). Cognitive load theory and instructional design: Recent developments. *Educational Psychologist*, 38(1), 1–4.
- Pinos Ullauri, L. A. (2024). *Leveraging latent variables to support learning* [Doctoral dissertation, Institut Mines-Télécom Nord Europe; KU Leuven] [NNT: 2024MTLD0013].
- Shute, V. J. (2008). Focus on formative feedback. *Review of educational research*, 78(1), 153–189.

van Merriënboer, J. J. G., & Ayres, P. (2005). Research on cognitive load theory and its design implications for e-learning. *Educational Technology Research and Development*, 53(3), 5–13.

Appendix A: Instructional Materials Metadata and expert review

Table S1: Sample Q-matrix for the simulated assessment (first 4 of 60 items shown). A value of 1 indicates that the item requires mastery of that skill; 0 indicates it does not. Full matrix follows the same structure across all 60 items and 5 skills, with 60% of items requiring one skill and 40% requiring two or three skills.

Item	Skill 1	Skill 2	Skill 3	Skill 4	Skill 5
1	1	0	0	0	0
2	0	1	1	0	0
3	0	0	1	0	0
4	1	0	0	1	0
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
60	0	0	1	0	1

All items exhibited positive discrimination estimates ($a_1 > 0$) and stable asymptote constraints ($g = 0, u = 1$), indicating no pathological item behavior under the fitted model.

Table S2: IRT item parameters for the 19 operational items. Note: in the estimation output, row labels 12–30 are internal IDs; operational Item i corresponds to internal ID $i + 11$ (e.g., internal ID 12 is operational Item 1). All estimated items show acceptable parameter values under the fitted model (here $g = 0$ and $u = 1$ for all items).

Operational Item	a_1	d	g	u
1	0.9723339	0.67592911	0	1
2	0.8138882	-0.03907757	0	1
3	0.6131102	-0.47164060	0	1
4	1.5157102	1.17200838	0	1
5	0.3200164	-0.34142756	0	1
6	0.9334762	0.43647420	0	1
7	1.8403050	0.93331301	0	1
8	1.0615326	0.15470139	0	1
9	2.3482633	0.45745660	0	1
10	1.5254729	0.51193608	0	1
11	1.1518942	0.37099326	0	1
12	1.6251323	0.41362895	0	1
13	0.8358662	0.65009197	0	1
14	1.4094174	0.80601187	0	1
15	1.3904804	-0.08549747	0	1
16	2.4513047	1.10707906	0	1
17	0.9463458	-0.05689655	0	1
18	1.4284152	0.45443303	0	1
19	0.5540262	-0.67610499	0	1

Table S3: DINA model fit indices and model size.

Quantity	Value
Number of cases (N)	720
Number of items (I)	19
Number of skill dimensions (K)	4
Number of skill classes	16
Number of parameters (total)	53
Item parameters	38
Skill distribution parameters	15
Log-likelihood	-8605.42
AIC	17317
BIC	17562

Appendix B: Greedy Heuristic for Content Selection

The following procedure describes the greedy heuristic used for selecting instructional resources for each student under skill coverage and time constraints.

1. Initialization.

- Let $V_i \leftarrow \emptyset$ be the set of selected resources for student i .
- Let $S \leftarrow S_i$ be the set of non-mastered skills for student i .
- Let $T \leftarrow 0$ be the cumulative viewing time.

2. Iterative selection. While $S \neq \emptyset$ and $T < T_{\max}$, repeat:

(a) Candidate identification. Identify candidate resources

$$C_i = \{j \in \{1, \dots, M\} : |S \cap C_j| > 0\},$$

where C_j denotes the set of skills addressed by resource j .

(b) Candidate scoring. For each $j \in C_i$, compute:

$$\text{Coverage}_j = |S \cap C_j|,$$

$$\text{Penalty}_j = \varepsilon L_j + p(1 - \delta_{D_j, P_i}),$$

$$F_{ij} = \text{Coverage}_j - \text{Penalty}_j.$$

(c) Selection. Select the resource

$$j^* = \arg \max_{j \in C_i} F_{ij}.$$

(d) State update.

- Update selected resources: $V_i \leftarrow V_i \cup \{j^*\}$.
- Remove covered skills: $S \leftarrow S \setminus C_{j^*}$.
- Update time: $T \leftarrow T + L_{j^*}$.

3. Output. Return the selected resource set V_i .

Table S4: **Expert–model agreement on resource selection and assigned duration.** GREEDY_LLM output compared against three independent instructor selections across ten representative student profiles.

Student	θ	Level	Gaps	Model		Instructor 1		Instructor 2		Instructor 3	
				Videos	Min	Videos	Min	Videos	Min	Videos	Min
S-01	−1.82	Low	S2,S3,S4	19_2, 20_4, 21_1	29	19_2, 20_4, 21_1	29	19_2, 20_4, 22_1	31	19_2, 20_4, 21_1	29
S-02	−0.94	Low	S2,S4	19_2, 20_1	20	19_2, 20_1	20	19_2, 20_1	20	19_3, 20_1	23
S-03	+0.11	Mid	S3,S4	20_4, 22_1	23	20_4, 22_1	23	20_4, 22_2	20	21_3, 22_1	24
S-04	−1.35	Low	S2,S3	19_2, 20_4	18	19_2, 20_4	18	19_2, 20_4	18	19_2, 21_3	19
S-05	+0.67	Mid	S2	19_3	11	19_3	11	19_3	11	19_3	11
S-06	+1.24	High	S4	22_3	15	22_3	15	22_1	13	22_3	15
S-07	+0.33	Mid	S2,S4	23_1	12	23_1	12	23_2	13	23_1	12
S-08	−0.58	Low	S1,S2,S4	22_4, 20_1	27	22_4, 20_1	27	22_4, 20_1	27	22_4, 21_1	24
S-09	+1.51	High	S2,S4	23_2	13	23_2	13	23_1	12	23_2	13
S-10	+0.89	High	S3,S4	24_1, 22_2	20	24_1, 22_2	20	24_1, 22_3	25	24_1, 22_2	20
<i>Resource identity agreement</i>				—		100%		70%		80%	
<i>Duration agreement (± 5 min)</i>				—		100%		80%		70%	
Mean across instructors				—		100%		75%		75%	

Note. Skill gap labels — S1: systems & trigonometric setup; S2: torque; S3: angular momentum; S4: principle of angular momentum. Mean resource identity agreement across instructors: $(100 + 70 + 80)/3 \approx 83\%$. Mean duration agreement: $(100 + 80 + 70)/3 \approx 83\%$. Video 22_4 covers Skills 1, 2, and 4 simultaneously; the model exploits this to minimize assigned resources for multi-gap students (e.g., S-08). Instructors 2 and 3 occasionally preferred difficulty-matched alternatives within the same skill, accounting for the observed divergence.

Appendix C: Gradient Descent for Content Selection

This procedure describes the continuous optimization approach used for content selection via gradient descent under multiple constraints.

1. **Initialization.** For each student $i = 1, \dots, N$ and resource $j = 1, \dots, M$, initialize assignment variables

$$x_{ij} \leftarrow 0,$$

where $x_{ij} \in [0, 1]$ represents the probability of assigning resource j to student i .

2. **Iterative optimization.** Repeat until convergence:
 - (a) **Gradient computation.** For each student i and resource j , compute:
 - *Skill coverage penalty:*

$$\nabla_{x_{ij}}^{\text{coverage}} = -2\lambda_{\text{coverage}} \max\left(0, (1 - S_{ik}) - \sum_{j'=1}^M x_{ij'} C_{j'k}\right) C_{jk}.$$

- *Watch-time constraint:*

$$\nabla_{x_{ij}}^{\text{time}} = 2\lambda_{\text{time}} \max\left(0, \sum_{j'=1}^M x_{ij'} L_{j'} - T_{\text{max}}\right) L_j.$$

- *Difficulty mismatch penalty:*

$$\nabla_{x_{ij}}^{\text{fallback}} = \lambda_{\text{fallback}} F_{ij}.$$

- *Coverage–efficiency tradeoff:*

$$\nabla_{x_{ij}}^{\text{utility}} = -\sum_{k=1}^K (1 - S_{ik}) C_{jk} + \varepsilon L_j.$$

Table S5: Distributional summaries of outcome and alignment metrics (by method).

Method	Metric	Mean	SD	Median	IQR	P10	P90	Min	Max
GRADIENTDESCENT_LLM	Utility	2.988	1.121	2.674	1.142	1.695	4.599	1.433	7.378
GRADIENTDESCENT_LLM	Assigned resources	7.454	5.345	5.000	6.000	2.000	15.000	1.000	36.000
GRADIENTDESCENT_LLM	Utility per resource	0.420	0.301	0.344	0.307	0.151	0.788	0.050	2.748
GRADIENTDESCENT_LLM	LLM similarity (mean)	0.665	0.192	0.721	0.242	0.349	0.858	0.192	0.905
GRADIENTDESCENT_LLM	Coverage rate	0.996	0.027	1.000	0.000	1.000	1.000	0.667	1.000
GREEDY_LLM	Utility	2.443	1.074	2.157	1.144	1.212	3.976	0.513	6.584
GREEDY_LLM	Assigned resources	3.901	3.292	3.000	3.000	1.000	8.000	1.000	19.000
GREEDY_LLM	Utility per resource	0.651	0.379	0.550	0.379	0.286	1.100	0.092	3.305
GREEDY_LLM	LLM similarity (mean)	0.626	0.195	0.687	0.283	0.297	0.848	0.110	0.878
GREEDY_LLM	Coverage rate	1.000	0.000	1.000	0.000	1.000	1.000	1.000	1.000

Table S6: Distributional summaries of constraint-related metrics (by method).

Method	Metric	Mean	SD	Median	IQR	P10	P90	Min	Max
GRADIENTDESCENT_LLM	Missing required skills	0.025	0.156	0.000	0.000	0.000	0.000	0.000	1.000
GRADIENTDESCENT_LLM	Extra covered skills	3.690	3.677	3.000	3.000	0.000	8.000	0.000	20.000
GREEDY_LLM	Missing required skills	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
GREEDY_LLM	Extra covered skills	2.194	2.528	2.000	2.000	0.000	6.000	0.000	17.000

(b) **Gradient aggregation.** Combine all gradient components:

$$\nabla x_{ij} = \nabla_{x_{ij}}^{\text{coverage}} + \nabla_{x_{ij}}^{\text{time}} + \nabla_{x_{ij}}^{\text{fallback}} + \nabla_{x_{ij}}^{\text{utility}}.$$

(c) **Variable update.** Update assignment probabilities:

$$x_{ij} \leftarrow x_{ij} - \eta \nabla x_{ij},$$

and project onto the interval $[0, 1]$.

(d) **Convergence check.** Stop if $\|\nabla x_{ij}\| < \delta$.

3. **Final assignment.** Set

$$x_{ij} = \begin{cases} 1, & \text{if } x_{ij} \geq 0.5, \\ 0, & \text{otherwise,} \end{cases}$$

and return

$$V_i = \{j : x_{ij} = 1\}.$$

Table S7: Ability-stratified performance profiles (means within ability quartiles).

Method	Ability Quartile	n	Satisfactory Rate	Utility	Missing Skills	Assigned Resources	LLM Similarity (Mean)
GRADIENTDESCENT_LLM	Q1	180	0.967	3.168	0.033	8.406	0.623
GRADIENTDESCENT_LLM	Q2	180	0.978	2.939	0.022	7.456	0.668
GRADIENTDESCENT_LLM	Q3	180	0.978	2.886	0.022	6.661	0.679
GRADIENTDESCENT_LLM	Q4	180	0.978	2.958	0.022	7.294	0.692
GREEDY_LLM	Q1	180	1.000	2.657	0.000	4.767	0.581
GREEDY_LLM	Q2	180	1.000	2.391	0.000	3.639	0.629
GREEDY_LLM	Q3	180	1.000	2.310	0.000	3.594	0.643
GREEDY_LLM	Q4	180	1.000	2.414	0.000	3.606	0.652

Table S8: Assignment diversity over the resource repository (usage concentration and effective support).

Method	Assigned resource instances	Unique resources used	Top-1 share	Top-5 share	Top-10 share	HHI	Effective resource count
GRADIENTDESCENT_LLM	5367	48	0.076	0.314	0.504	0.035	28.620
GREEDY_LLM	2809	48	0.131	0.416	0.585	0.048	20.923

Table S9: Distribution of covered skills aggregated over student assignments.

Method	Covered skill	Unique skills	Top-1 share	Top-5 share	HHI	Effective skill
GRADIENTDESCENT_LLM	4952	4	0.373	1.000	0.284	3.527
GREEDY_LLM	3021	4	0.377	1.000	0.285	3.511

Table S10: Stability and overlap across models for paired students (distribution of Jaccard similarity).

Metric	Mean	SD	Median	IQR	P10	P90	Min	Max
Jaccard similarity (assigned resources)	0.245	0.215	0.200	0.286	0.000	0.500	0.000	1.000
Jaccard similarity (covered skills)	0.777	0.258	1.000	0.500	0.333	1.000	0.000	1.000

Table S11: Paired trade-off patterns comparing GRADIENTDESCENT_LLM to GREEDY_LLM (same students).

Trade-off pattern	Paired students	Proportion
GradientDescent utility gain w/ more resources	316	87.8%
GradientDescent fewer/equal resources w/o utility gain	44	12.2%
GradientDescent more resources w/o utility gain	0	0.0%

Table S12: Coverage conditions under paired comparison (same students).

Coverage condition (paired)	Paired students	Proportion
GradientDescent has any missing required skills	19	5.3%
GradientDescent missing worsens vs Greedy (paired)	19	5.3%
GradientDescent and Greedy both fully cover required skills	341	94.7%
GradientDescent has exact coverage parity with Greedy	341	94.7%

Table S13: Coverage-first dominance outcomes for paired students.

Outcome	Students	Proportion
Non-comparable	331	91.9%
Exact tie	29	8.1%