

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Treating Distillation as Pedagogy: Simplicity and Scaffolding in Large Language Models

Permalink

<https://escholarship.org/uc/item/67g410sz>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Chen, Jian

Huang, Fuhua

Liu, Guojie

et al.

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Treating Distillation as Pedagogy: Simplicity and Scaffolding in Large Language Models

Jian Chen^{*,1,2}, Fuhua Huang², Guojie Liu², Kun Shen², Linwei Liu¹ & Xu Lu¹

¹Ningxia Transport Science Research Institute Co., Ltd., Yinchuan, China

²Ningxia Highway Engineering Quality Inspection Center (Co., Ltd.), Yinchuan, China

Abstract

Human pedagogy relies on scaffolding to tailor information complexity to a learner’s capacity, yet current AI knowledge distillation often subjects resource-constrained student models to computational cognitive overload via unstructured, verbose reasoning. To investigate the alignment between human learning principles and machine learning dynamics, we propose a unified computational framework that treats distillation as a controlled cognitive experiment. We manipulated the *granularity* and *sequence* of the pedagogical signal across two distinct tasks. Our results demonstrate two key phenomena paralleling human cognition: a simplicity effect, where concise elucidations significantly outperform complex expert rationales by minimizing extraneous cognitive load, and a scaffolding effect, where a simple-to-complex curriculum proves essential for convergence in reasoning tasks. These findings provide computational evidence that less is more in the context of knowledge distillation for resource-constrained models, suggesting that effective alignment in student-teacher paradigms requires a shift from maximizing information quantity to optimizing pedagogical quality.

Keywords: Cognitive Load Theory; Knowledge Distillation; Large Language Models; Scaffolding; Curriculum Learning

Introduction

Human pedagogy is fundamentally constrained by the information processing limits of the learner. According to Cognitive Load Theory (CLT), efficient learning occurs only when the total cognitive load—comprising intrinsic, extraneous, and germane loads—remains within the learner’s working memory capacity (Sweller, 1988; Paas & Van Merriënboer, 2020). Consequently, expert human teachers instinctively employ two critical strategies to prevent cognitive overload in novices: simplification (reducing extraneous details to maximize signal-to-noise ratio) and scaffolding (structuring learning material from simple to complex) (Wood et al., 1976; Vygotsky, 1978). Experts who fail to adapt their instruction in this way often suffer from the expert blind spot, baffling novices with explanations that are accurate but overly complex (Nathan & Petrosino, 2003).

In the domain of artificial intelligence, specifically in Knowledge Distillation (KD) for Large Language Models (LLMs) (Hinton et al., 2015), a striking parallel has emerged. The standard paradigm involves a high-capacity Teacher model (e.g., GPT-5 (Achiam et al., 2023)) transferring knowl-

edge to a resource-constrained Student model (e.g., 4B parameters small LLMs (Yang et al., 2025)). Recent advances have shifted towards reasoning distillation, where students are trained to mimic the Teacher’s verbose Chain-of-Thought (CoT) traces (Wei et al., 2022; Guo et al., 2025). However, this approach often adopts a brute-force pedagogy: students are bombarded with complex, multi-step reasoning chains in a random order, regardless of their current representational capacity.

We hypothesize that machine learners, akin to human novices, are subject to a form of computational cognitive overload (see Figure 1). We argue that the prevailing more is better assumption in AI—where richer, longer, and more complex teacher outputs are always presumed superior—is misaligned when applied to novice learners with limited channel capacity. This aligns with recent findings on the ‘simplicity bias’ of LLMs, where excessively long reasoning chains have been shown to degrade performance due to error accumulation (Y. Wu et al., 2025). When the complexity of the teacher’s reasoning exceeds the student’s bandwidth, the detailed rationale transforms from a helpful signal into confusing noise, hindering convergence and generalization.

To empirically test this hypothesis, we propose a unified computational framework that treats the Teacher-Student distillation process as a controlled cognitive experiment. We investigate two core research questions:

- **RQ1 (Content):** Does the *simplicity principle* hold for novice machines? i.e., Is a concise elucidation more effective than a complex expert rationale for student domain adaptation?
- **RQ2 (Sequence):** Does the *scaffolding principle* hold for machines? i.e., Is a simple-to-complex curriculum essential for acquiring reasoning capabilities?

Through controlled experiments on two distinct tasks—domain concept acquisition and narrative reasoning—we provide computational evidence that aligns closely with human learning theories. Our results demonstrate that optimizing the *pedagogical quality* of training data yields superior performance compared to maximizing information quantity, suggesting that future AI systems require pedagogical intelligence to bridge the gap between expert models and novice learners.

*Corresponding author: jchen.ai@foxmail.com.

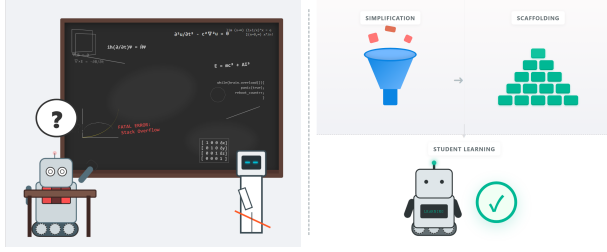


Figure 1: **The expert blind spot vs. cognitive-aligned distillation.** (Left) Current paradigms subject student models to *computational cognitive overload* by exposing them to complex, unstructured expert reasoning, leading to optimization failures. (Right) Our framework treats distillation as a pedagogical process, employing **simplification** and **scaffolding** to align with the student’s learning capacity.

Related Works

Our work sits at the intersection of cognitive science, pedagogy, and machine learning. We review relevant literature in human learning theories and their computational counterparts in knowledge distillation.

Cognitive Load Theory and Pedagogical Scaffolding

Human pedagogy is fundamentally constrained by the limited capacity of working memory. Cognitive Load Theory (CLT) posits that learning is hampered when the total cognitive load exceeds this capacity (Sweller, 1988; Paas & Van Merriënboer, 2020). Sweller distinguishes between *intrinsic load* and *extraneous load*. Effective instruction must minimize extraneous load to facilitate schema acquisition.

Parallel to this, the concept of scaffolding, rooted in Vygotsky’s Zone of Proximal Development (ZPD) (Vygotsky, 1978), suggests that learners require support structures that are dynamically adjusted to their current competence level (Wood et al., 1976). Instruction that is too complex leads to confusion, while instruction that is too simple leads to stagnation. Our study operationalizes these human-centric concepts into computational constraints for machine learners.

Knowledge Distillation and the Expert Blind Spot

Knowledge Distillation (KD) aims to transfer knowledge from a Teacher to a Student. While early methods focused on logit-based mimicry (Hinton et al., 2015), the advent of Large Language Models (LLMs) (Achiam et al., 2023) has shifted the paradigm towards transferring reasoning capabilities via Chain-of-Thought (CoT) traces (Guo et al., 2025).

However, a critical misalignment often occurs: expert models tend to generate highly detailed and verbose reasoning paths. In educational psychology, this phenomenon is known as the expert blind spot (Nathan & Petrosino, 2003), where experts fail to anticipate the difficulties novices face with complex intermediate steps. In the context of LLMs, this mirrors the misalignment where teacher models optimize for generation probability rather than instructional utility (Sonkar et

al., 2024). We hypothesize that current CoT-based distillation methods inadvertently introduce high *extraneous load* for small student models, causing what we term computational cognitive overload.

Curriculum Learning in Artificial Intelligence

Inspired by human learning, Curriculum Learning (CL) proposes training machine learning models on samples organized from easy to difficult (Bengio et al., 2009). While CL has shown promise in improving convergence, defining “difficulty” remains an open challenge.

Existing approaches often rely on heuristic metrics or model-based uncertainty (Soviany et al., 2022; Wang et al., 2021). However, recent work specifically targeting LLM reasoning suggests that adaptive difficulty adjustment is crucial for acquiring complex logic (E. Zhang et al., 2025). Our work diverges by proposing a *cognitive-inspired* difficulty metric: the length of the reasoning trace required to solve a problem. This aligns with the cognitive view that problems requiring more working memory retention are inherently more cognitively demanding, necessitating a strict scaffolding strategy.

Computational Framework

To empirically test the applicability of human pedagogical principles—specifically *simplification* and *scaffolding*—to the domain of knowledge distillation in LLMs, we propose a unified computational simulation framework as shown in Figure 2.

This framework establishes a controlled Teacher-Student environment, allowing us to isolate and manipulate two critical pedagogical variables: the *complexity of the content* (Experiment 1) and the *sequence of instruction* (Experiment 2). We present these experimental designs as novel probing mechanisms to evaluate cognitive load constraints in machine learners.

Teacher-Student Paradigm

We formalize the learning process as a knowledge transfer task from a high-capacity expert model ($\mathcal{M}_T$) to a resource-constrained novice model ($\mathcal{M}_S$).

The Teacher ($\mathcal{M}_T$): We employ a large-scale foundation model as the pedagogical agent. This model possesses strong emergent reasoning capabilities and serves as the source of knowledge generation.

The Student ($\mathcal{M}_S$): We employ a compact architecture to simulate a learner with limited cognitive capacity. The student’s objective is to minimize the loss $\mathcal{L}$ on a downstream task, conditioned on the pedagogical signal provided by the teacher.

Experiment 1: Content Modulation

To verify the hypothesis that simple explanations facilitate learning in novices better than complex reasoning, we design a contrastive experiment focusing on the granularity and signal-to-noise ratio of the generated knowledge.

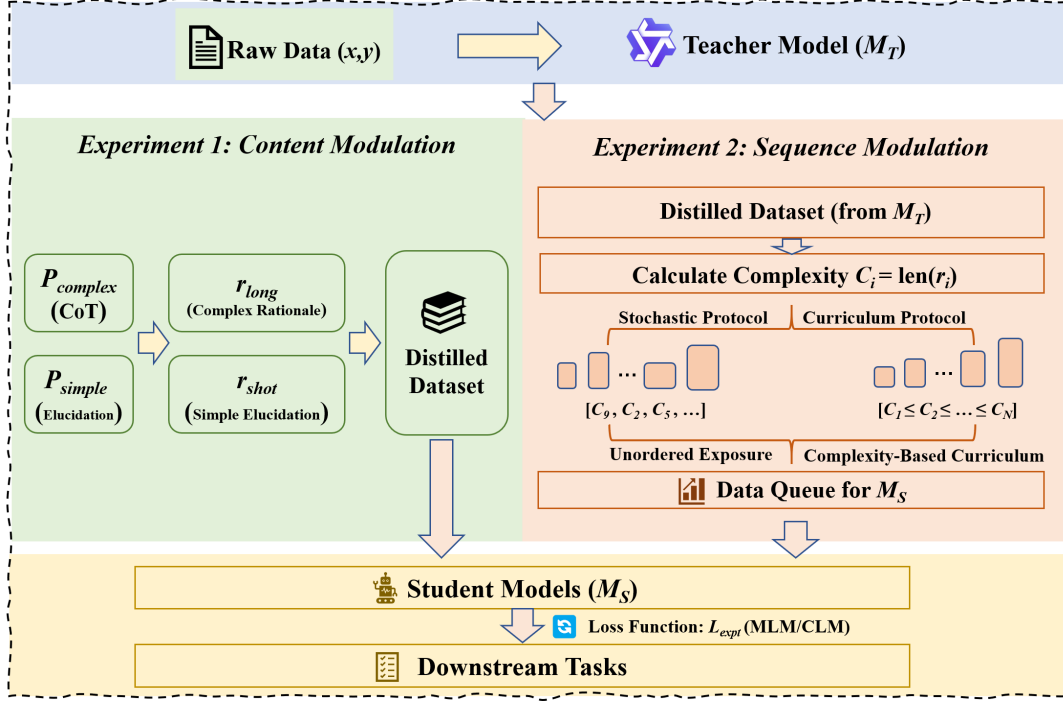


Figure 2: **The unified computational framework for cognitive-aligned distillation.** We treat distillation as a controlled experiment manipulating two variables: (1) Content Granularity, contrasting complex rationales against simple elucidations, and (2) Instructional Sequence, comparing stochastic exposure against a complexity-based curriculum.

Task Definition: The goal is domain adaptation. The student model must update its internal representations to understand specialized concepts before performing extraction tasks (Z. Zhang et al., 2024). Given an input text x and a target label y , the Teacher M_T generates auxiliary knowledge $\mathcal{K}$ based on a specific instructional prompt P : $\mathcal{K} \sim P_{M_T}(\cdot|x, y, P)$. We design two conditions (detailed in Table 1):

- **Condition A (Complex Rationale):** The prompt $P_{complex}$ instructs the teacher to generate a verbose, multi-step Chain-of-Thought (CoT) trace, simulating an expert’s full cognitive process with functional analysis.
- **Condition B (Simple Elucidation):** The prompt P_{simple} instructs the teacher to generate a concise, direct explanation that explicitly links the input span to the target category, filtering out extraneous inferential noise.

To simulate the student studying this material, we employ an incremental pre-training objective. The student M_S is trained to reconstruct the pedagogical signal $\mathcal{K}$ via Masked Language Modeling (MLM). Let $\tilde{\mathcal{K}}$ be the corrupted version of $\mathcal{K}$ with random masking. The optimization objective is:

$$\mathcal{L}_{\text{exp1}} = - \sum_{t \in \text{mask}} \log P(k_t | \tilde{\mathcal{K}}; \theta_S) \quad (1)$$

This formulation allows us to measure whether the information density of $\mathcal{K}$ aligns with the student’s processing capacity. A lower loss on the downstream task implies successful internalization of the pedagogical signal.

Common Prompt Instruction (Invariant)	
<i>"You are a world-class requirements analysis expert... Explain why a given text fragment is labeled as a specific entity type... Your explanation must be logical and well-grounded."</i>	
Condition	Generative Mechanism & Signal Characteristics
Complex Rationale	Mechanism: Qwen3 Thinking Mode (ON) $\hookrightarrow$ Triggers internal chain-of-thought, recursive analysis, and self-correction before outputting the final explanation.
(High Load)	Signal: Highly verbose, captures the unstructured stream of consciousness of an expert, maximizing intrinsic and extraneous load.
Simple Elucidation	Mechanism: Qwen3 Standard Mode (OFF) $\hookrightarrow$ Direct generation based on probability distribution without latent reasoning steps.
(Low Load)	Signal: Concise and result-oriented. It filters out the inferential search process, providing a high signal-to-noise ratio explanation.

Table 1: **Contrast of pedagogical signal generation.** While the input prompt remains constant, we manipulate the teacher’s cognitive depth. The *complex* condition activates the model’s thinking mode, producing exhaustive traces. The *simple* condition uses standard generation, yielding direct elucidations.

Experiment 2: Sequence Modulation

To verify the hypothesis that scaffolding via curriculum sequencing improves convergence, we design an experiment that manipulates the order in which data is presented to the student.

Task Definition: The goal is complex reasoning acquisition. The student must learn to generate a reasoning path r followed by a correct classification y , modeled as an autoregressive sequence generation task. We quantify the cognitive load of each training sample i using the length of the Teacher’s reasoning trace as a proxy. Let r_i be the trace generated by $\mathcal{M}_T$. We define the complexity score as $C_i = \text{len}(r_i)$, based on the assumption that longer reasoning paths require higher retention of intermediate states in the student’s working memory.

We train the Student $\mathcal{M}_S$ on the distilled dataset $\mathcal{D} = \{(x_i, r_i, y_i)\}_{i=1}^N$. The training objective is the standard Causal Language Modeling (CLM) loss:

$$\mathcal{L}_{\text{exp2}} = - \sum_{j=1}^L \log P(z_j | x_i, z_{<j}; \theta_S) \quad (2)$$

where z represents the concatenated sequence of reasoning r_i and label y_i . We compare two protocols:

- **Protocol A (Stochastic / No Scaffolding):** Samples are presented in a random order. This simulates an unstructured learning environment where high-variance gradients from complex samples may destabilize early training.
- **Protocol B (Curriculum / Scaffolding):** Samples are sorted by their complexity score C_i in ascending order:

$$C_{\pi(1)} \leq C_{\pi(2)} \leq \dots \leq C_{\pi(N)} \quad (3)$$

The student is exposed to samples with short, direct reasoning paths before progressing to those requiring long-context inference, allowing the model to establish stable basins of attraction in the optimization landscape.

Experiments

Experimental Setup

To evaluate our hypotheses under varying cognitive demands, we employ two distinct experimental environments covering different tasks and data modalities.

Environment 1: Domain Concept Acquisition. To test the simplicity effect, we utilize the Architecture, Engineering, and Construction (AEC) benchmark dataset (Zhou et al., 2022). This dataset comprises 611 specialized sentences with 4,336 annotated entities. The task represents a high-difficulty domain adaptation scenario where a novice model must bridge a significant linguistic gap to understand specialized terminology.

Environment 2: Narrative Reasoning. To test the scaffolding effect, we utilize a distilled multimodal variant of the OSHA benchmark (Goh & Ubeynarayana, 2017). Starting with the full corpus of 1,000 accident narratives, we applied a standard 80/20 random split, allocating 800 samples for training and 200 for testing. To simulate a multimodal reasoning

scenario, we reframed the original text-only task by augmenting each narrative with a synthesized image generated via Qwen-Image(C. Wu et al., 2025). Subsequently, we implemented a rigorous data sanitization protocol on the training split to ensure pedagogical quality. This filtering process reduced the initial 800 training samples to a high-fidelity subset of 477 verified examples. Characterized by severe class imbalance and long-context dependencies, this verified subset provides a challenging testbed for evaluating the stability of curriculum learning.

Model Architectures: The Pedagogical Agents

We instantiate the Teacher-Student paradigm using models of varying capacities to simulate the Expert-Novice gap.

For Experiment 1 (Content Modulation):

- **Teacher ($\mathcal{M}_T$):** We employ Qwen3-8B (Yang et al., 2025), a general-purpose model. In this setup, the teacher is responsible for generating elucidations to bridge the AEC domain gap.
- **Student ($\mathcal{M}_S$):** We employ RoBERTa-wwm-ext (Cui et al., 2020), a RoBERTa-based (Liu et al., 2019) encoder model. Representing a learner with limited semantic capacity, it relies on the teacher’s signals to update its internal representations before fine-tuning on the downstream NER task.

For Experiment 2 (Sequence Modulation):

- **Teacher ($\mathcal{M}_T$):** We employ Qwen3-VL-30B-A3B-Thinking (Bai et al., 2025), a large-scale vision-language model optimized for reasoning. This teacher generates the CoT traces that serve as the learning material.
- **Student ($\mathcal{M}_S$):** We employ Qwen3-VL-4B-Instruct (Bai et al., 2025), a lightweight multimodal model. As a novice generative model, it attempts to emulate the teacher’s reasoning path via autoregressive distillation.

Implementation Details

Experiment 1 Protocol: We compare two distinct knowledge generation strategies. The *complex condition* uses prompts that trigger verbose, role-based analysis from the Teacher, while the *simple condition* uses prompts that elicit direct, context-bound elucidations.

Experiment 2 Protocol: We compare two training sequences on the distilled data. The *stochastic protocol* feeds samples to the student in random order. The *curriculum protocol* sorts samples based on the length of the teacher’s CoT trace, ensuring the student masters short reasoning paths before attempting long-context inferences. All experiments were conducted on a NVIDIA RTX 4090 GPU using the PyTorch (Paszke et al., 2019) and Transformers (Wolf et al., 2020) framework.

Results and Cognitive Analysis

Experiment 1: We compared the performance of the student model trained with simple elucidations against the one trained

Table 2: Main results. **Panel A** demonstrates the simplicity effect (Exp 1), comparing pedagogical signal granularity. **Panel B** demonstrates the scaffolding effect (Exp 2), comparing learning sequences. **Bolded** values indicate the best performance.

Models	Method / Protocol	Precision	Recall	Macro-F1
Panel A: Experiment 1 - Content Modulation				
BERT-base-Chinese (Devlin et al., 2019)	-	77.54	71.02	72.93
RoBERTa-wwm-ext (Cui et al., 2020)	-	74.62	68.36	71.19
ARCBERT (Z. Zheng et al., 2022)	-	77.31	75.26	75.75
Qwen3-8B-SFT (Y. Zheng et al., 2024)	-	75.48	75.39	75.23
Qwen-AEC-8B-NER (Chen et al., 2026)	-	75.69	76.33	75.37
 Ours	Complex Rationale	76.46	77.67	76.86
	Simple Elucidation	77.51	77.30	77.20
Panel B: Experiment 2 - Sequence Modulation				
BERT-base-uncased (Devlin et al., 2019)	-	66.29	63.48	62.45
RoBERTa-base (Liu et al., 2019)	-	57.49	63.27	58.73
Qwen3-VL-4B-Instruct (Bai et al., 2025)	-	63.86	62.50	58.35
Qwen3-VL-4B (SFT w/o CoT) (Y. Zheng et al., 2024)	-	62.64	67.85	63.47
INSTRUCTOR-CIT (Shuang et al., 2024)	-	66.09	63.89	64.22
 Ours	Protocol A (Stochastic)	63.64	69.23	61.84
	Protocol B (Curriculum)	66.76	72.20	65.23

with complex rationales. Table 2 summarizes the results on the AEC domain adaptation task.

The student trained on simple explanations achieves a superior Macro-F1 score compared to the complex condition. While the complex rationale yields a slightly higher recall, it suffers from a significant drop in precision. These results provide computational evidence for the detrimental effects of *extraneous cognitive load*. The complex rationale condition, mimicking an expert’s verbose thought process, casts a wider semantic net, hence the marginal gain in recall. However, this verbosity introduces inferential noise that exceeds the student model’s processing capacity, leading to hallucinations and a drop in precision. In contrast, the simple elucidation acts as a high-fidelity information bottleneck, filtering out irrelevant causal links and maximizing the signal-to-noise ratio. This confirms that for novice learners, concise guidance is often more effective than exhaustive reasoning.

Experiment 2: We evaluated the impact of instructional sequencing by comparing the curriculum protocol (Protocol B) against the stochastic protocol (Protocol A) and a baseline without reasoning distillation. Table 2 presents the results on the imbalanced OSHA narrative classification task.

The most striking finding is that simply adding Chain-of-Thought reasoning in a random order actually *degrades* performance compared to the baseline SFT. However, when the exact same data is organized via a simple-to-complex curriculum, the performance leaps to a state-of-the-art Macro-F1 of 65.23%.

This result offers a validation of *cognitive load theory* in a machine learning context. Protocol A demonstrates computational cognitive overload. When a novice model is prematurely exposed to long, complex reasoning traces without prior structure, the high-variance gradients destabilize the learning

process, making the extra information detrimental rather than helpful. Protocol B simulates the zone of proximal development. By mastering short, direct reasoning paths first, the student establishes stable feature representations. This foundational knowledge serves as a scaffold, allowing the model to subsequently accommodate and benefit from complex, long-context inferences.

Case Study: To provide concrete insight into the mechanisms of *computational cognitive overload*, we present qualitative examples from both experiments. Figure 3 illustrates the simplicity effect in the AEC task: while the expert’s verbose reasoning introduces semantic ambiguity, the concise elucidation acts as a high-fidelity signal, ensuring correct alignment.

Complementarily, Figure 4 demonstrates the scaffolding effect in the multimodal OSHA task. Unlike the stochastic protocol, which subjects the learner to erratic shifts in visual and reasoning complexity, the curriculum protocol stabilizes optimization by ensuring the model masters simple, direct mappings before attempting complex causal reasoning.

Conclusion

This study bridges educational psychology and AI by demonstrating that machine learners, akin to human novices, are critically sensitive to pedagogical quality. We empirically validated two cognitive principles in knowledge distillation: a simplicity effect and a scaffolding effect, suggesting that effective AI alignment requires optimizing pedagogical intelligence—dynamically adapting the granularity and sequence of information to the learner’s specific capacity rather than merely scaling data quantity. Future work will extend this framework to agent and embodied intelligence, further exploring the convergence between biological and artificial cognition.

Case Study 1: The Impact of Pedagogical Signal Granularity

Input Text:

“The head loss [Target 1] of the pressure reducing valve [Target 2] should be less than the design static and dynamic pressure difference...”

Condition A: Complex Rationale (Teacher with Thinking Mode)

Teacher's Thought Trace: *<think> I need to analyze 'pressure reducing valve'. Definition check: An 'Object' ('obj') is a parent node of attributes. Here, it is associated with 'head loss'. It is also a mechanical device used for pressure control. Is it 'Equipment' or 'Object'? The text emphasizes its attributes (head loss, pressure diff). It fits the 'Object' definition as a property carrier. Similarly, 'head loss' is a performance parameter... [Verbose derivation follows] </think>*

→ **Teacher Output:** “Analyzing the semantic structure: 'Pressure reducing valve' functions as a physical mechanical device possessing performance metrics...”

→ **Student Prediction:** **sobj** ✗ (Incorrect - The description mechanical device triggered a domain mismatch error)

Condition B: Simple Elucidation (Teacher with Standard Mode)

Teacher's Direct Output: 1. “pressure reducing valve”: It is the parent entity containing attributes. **Label:** ‘obj’ (Object).

2. “head loss”: It is the specific attribute associated with the valve. **Label:** ‘prop’ (Property).

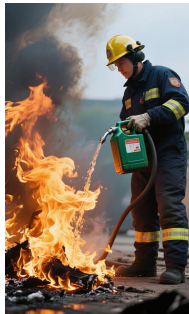
→ **Student Prediction:** **obj, prop** ✓ (Correct Alignment)

Figure 3: Qualitative analysis of pedagogical signals.

Case Study 2: The Impact of Instructional Sequence

Protocol A: Stochastic (Random)

High Variance: Abrupt shifts in complexity.



[Step t] **Complex Sample:** “...employee burned when leaked gasoline ignites...”

→ **Teacher (High Load):** “Analyze image & text: gasoline leaking → ignition → fire. Rule out ‘electrocution’, ‘falls’. Match category: *fires and explosion*.”



[Step $t + 1$] **Simple Sample:** “...employee dies after fall from roof...”

→ **Teacher (Low Load):** “Identify key phrase: ‘fall from roof’. No complex inference needed. Direct match category: *falls*.”

→ **Result:** *Unstable optimization landscape.*

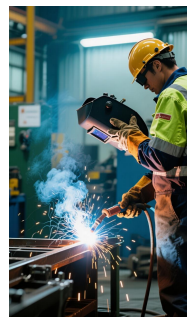
Protocol B: Curriculum (Sorted)

Scaffolding: Foundation on simple examples.



[Step t] **Simple Sample** ($C_t \approx L$): “...employee sustains electric shock...”

→ **Teacher:** “Check description: ‘electric shock’. Direct mapping found. Match category: *electrocution*.”



[Step $t + 1$] **Simple Sample** ($C_{t+1} \approx L$): “...electrocuted while arc welding...”

→ **Teacher:** “Check description: ‘electrocuted’. Consistent pattern observed. Match category: *electrocution*.”

→ **Result:** *Stable convergence & alignment.*

Figure 4: Qualitative comparison of training sequences

Acknowledgments

This study was supported by the Transportation Power Special Fund Project of the Ningxia Hui Autonomous Region (No. 202500210) and the Ningxia Natural Science Foundation (No. 2026AAC020063).

References

- Sweller, J. (1988). Cognitive load during problem solving: Effects on learning. *Cognitive science*, 12(2), 257–285.
- Paas, F., & Van Merriënboer, J. J. (2020). Cognitive-load theory: Methods to manage working memory load in the learning of complex tasks. *Current Directions in Psychological Science*, 29(4), 394–398.
- Wood, D., Bruner, J. S., & Ross, G. (1976). The role of tutoring in problem solving. *Journal of child psychology and psychiatry*, 17(2), 89–100.
- Vygotsky, L. S. (1978). *Mind in society: The development of higher psychological processes* (Vol. 86). Harvard university press.
- Nathan, M. J., & Petrosino, A. (2003). Expert blind spot among preservice teachers. *American educational research journal*, 40(4), 905–928.
- Hinton, G., Vinyals, O., & Dean, J. (2015). Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*.
- Achiam, J., Adler, S., Agarwal, S., Ahmad, L., Akkaya, I., Aleman, F. L., Almeida, D., Altenschmidt, J., Altman, S., Anadkat, S., et al. (2023). Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Yang, A., Li, A., Yang, B., Zhang, B., Hui, B., Zheng, B., Yu, B., Gao, C., Huang, C., Lv, C., et al. (2025). Qwen3 technical report. *arXiv preprint arXiv:2505.09388*.
- Wei, J., Wang, X., Schuurmans, D., Bosma, M., Xia, F., Chi, E., Le, Q. V., Zhou, D., et al. (2022). Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35, 24824–24837.
- Guo, D., Yang, D., Zhang, H., Song, J., Zhang, R., Xu, R., Zhu, Q., Ma, S., Wang, P., Bi, X., et al. (2025). Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*.
- Wu, Y., Wang, Y., Ye, Z., Du, T., Jegelka, S., & Wang, Y. (2025). When more is less: Understanding chain-of-thought length in llms. *arXiv preprint arXiv:2502.07266*.
- Sonkar, S., Ni, K., Chaudhary, S., & Baraniuk, R. (2024). Pedagogical alignment of large language models. *Findings of the Association for Computational Linguistics: EMNLP 2024*, 13641–13650.
- Bengio, Y., Louradour, J., Collobert, R., & Weston, J. (2009). Curriculum learning. *Proceedings of the 26th annual international conference on machine learning*, 41–48.
- Soviany, P., Ionescu, R. T., Rota, P., & Sebe, N. (2022). Curriculum learning: A survey. *International Journal of Computer Vision*, 130(6), 1526–1565.
- Wang, X., Chen, Y., & Zhu, W. (2021). A survey on curriculum learning. *IEEE transactions on pattern analysis and machine intelligence*, 44(9), 4555–4576.
- Zhang, E., Yan, X., Lin, W., Zhang, T., & Qianchun, L. (2025). Learning like humans: Advancing llm reasoning capabilities via adaptive difficulty curriculum learning and expert-guided self-reformulation. *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*, 6630–6644.
- Zhang, Z., Lee, S. Y. M., Wu, J., Zhang, D., Li, S., Cambria, E., & Zhou, G. (2024). Cross-domain ner with generated task-oriented knowledge: An empirical study from information density perspective. *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, 1595–1609.
- Zhou, Y.-C., Zheng, Z., Lin, J.-R., & Lu, X.-Z. (2022). Integrating nlp and context-free grammar for complex rule interpretation towards automated compliance checking. *Computers in Industry*, 142, 103746.
- Goh, Y. M., & Ubeynarayana, C. (2017). Construction accident narrative classification: An evaluation of text mining techniques. *Accident Analysis & Prevention*, 108, 122–130.
- Wu, C., Li, J., Zhou, J., Lin, J., Gao, K., Yan, K., Yin, S.-m., Bai, S., Xu, X., Chen, Y., et al. (2025). Qwen-image technical report. *arXiv preprint arXiv:2508.02324*.
- Cui, Y., Che, W., Liu, T., Qin, B., Wang, S., & Hu, G. (2020). Revisiting pre-trained models for chinese natural language processing. *arXiv preprint arXiv:2004.13922*.
- Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., & Stoyanov, V. (2019). Roberta: A robustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692*.
- Bai, S., Cai, Y., Chen, R., Chen, K., Chen, X., Cheng, Z., Deng, L., Ding, W., Gao, C., Ge, C., Ge, W., Guo, Z., Huang, Q., Huang, J., Huang, F., Hui, B., Jiang, S., Li, Z., Li, M., . . . Zhu, K. (2025). Qwen3-vl technical report. <https://arxiv.org/abs/2511.21631>
- Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., et al. (2019). Pytorch: An imperative style, high-performance deep learning library. *Advances in neural information processing systems*, 32.
- Wolf, T., Debut, L., Sanh, V., Chaumond, J., Delangue, C., Moi, A., Cistac, P., Rault, T., Louf, R., Funtowicz, M., et al. (2020). Transformers: State-of-the-art natural language processing. *Proceedings of the 2020 conference on empirical methods in natural language processing: system demonstrations*, 38–45.
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). Bert: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of the 2019 conference of the North American chapter of the association for computational linguistics: human language technologies, volume 1 (long and short papers)*, 4171–4186.

- Zheng, Z., Lu, X.-Z., Chen, K.-Y., Zhou, Y.-C., & Lin, J.-R. (2022). Pretrained domain-specific language model for natural language processing tasks in the aec domain. *Computers in Industry*, 142, 103733.
- Zheng, Y., Zhang, R., Zhang, J., Ye, Y., Luo, Z., Feng, Z., & Ma, Y. (2024). Llamafactory: Unified efficient fine-tuning of 100+ language models. *arXiv preprint arXiv:2403.13372*.
- Chen, J., Tian, J., Zheng, T., Hui, Y., & Li, Z. (2026). Qwen-aec: Instruction-tuning large language models for high-performance building code analysis. *Advanced Engineering Informatics*, 71, 104268.
- Shuang, Q., Liu, X., Wang, Z., & Xu, X. (2024). Automatically categorizing construction accident narratives using the deep-learning model with a class-imbalance treatment technique. *Journal of Construction Engineering and Management*, 150(9), 04024107.