

Highlights

A Theory-Guided LLM Pedagogical Agent for STEM+C Scaffolding Without Over-Reliance

Clayton Cohn, Surya Rayala, Siyuan Guo, Hanchen David Wang, Naveeduddin Mohammed, Umesh Timalisina, Shruti Jain, Ryan Li, Angela Eeds, Menton Deweese, Pamela J. Osborn Popp, Rebekah Stanton, Shakeera Walker, Ashwin T S, Meiyi Ma, Gautam Biswas

- Presents Copa, a theory-guided, agentic, multi-agent LLM pedagogical agent for STEM+C learning.
- Integrates multimodal learner evidence from student-agent interactions and environment logs.
- Reveals how Copa’s scaffolding evolves with students’ progress.
- Shows that students increasingly articulated conceptual understanding and confident reasoning as their task mastery developed.
- Demonstrates that Copa supported learning without fostering dependence, while keeping feedback traceable to student interaction data.

A Theory-Guided LLM Pedagogical Agent for STEM+C Scaffolding Without Over-Reliance

Clayton Cohn^a, Surya Rayala^a, Siyuan Guo^a, Hanchen David Wang^a, Naveeduddin Mohammed^a, Umesh Timalisina^a, Shruti Jain^a, Ryan Li^a, Angela Eeds^b, Menton Deweese^b, Pamela J. Osborn Popp^b, Rebekah Stanton^b, Shakeera Walker^b, Ashwin T S^a, Meiyi Ma^a and Gautam Biswas^a

^aCollege of Connected Computing, Vanderbilt University, Nashville, Tennessee 37240,

^bThe School for Science and Math at Vanderbilt, Vanderbilt University, Nashville, Tennessee 37240,

ARTICLE INFO

Keywords:

pedagogical agents
adaptive scaffolding
multi-agent
agentic
LLMs
interpretability
multimodal learning analytics
K-12 STEM+C

ABSTRACT

LLM pedagogical agents are proliferating, yet recent findings have raised questions about their adherence to established theories of learning and, by extension, their educational value. Concerns regarding cognitive offloading, over-reliance, and “gaming” behaviors persist and remain largely unaddressed. In response, we developed *Copa*, an agentic, multi-agent, multimodal Collaborative Peer Agent for STEM+C learning. *Copa* is built on top of the Evidence-Decision-Feedback (EDF) framework, grounding its interactions in Social Cognitive Theory and Social Constructivism and promoting sense-making through adaptive, dialogic support rather than answer-seeking. In an authentic high school computational-modeling study ($n = 33$ dyads), we demonstrate that *Copa* (1) supports students’ confidence building and ability to verbalize conceptual understanding without causing dependence; and (2) provides adaptive feedback personalized to learners that is interpretable with respect to students’ multimodal input data. These findings position theory-guided, multimodal LLM agents as a promising path toward classroom AI integration that amplifies students’ reasoning rather than replacing it.

1. Introduction

Large language model (LLM) pedagogical agents can support dialogic learning grounded in theories of social knowledge construction, including Social Cognitive Theory (SCT; e.g., self-efficacy and self-regulation) (Bandura, 2001) and Social Constructivism (e.g., Zone of Proximal Development [ZPD]) (Vygotsky, 1978). However, many LLM pedagogical agents, especially those external to artificial intelligence in education (AIED) venues (Zha, Liu, Zheng, Xu, Yu, Gong and Xu, 2025; Wang, Zhan, Lian, Hu, Yuan, Zhang, Xie and Xiong, 2025; Li, Xie, Chakravarty and Lee, 2024), do not exhibit the explicit theoretical grounding that characterized earlier intelligent tutoring systems (ITS), prompting calls to re-anchor these agents in learning theory (Lajoie and Li, 2023; Stamper, Xiao and Hou, 2024).

Insufficient attention to pedagogical theories in LLM agent design introduces significant risks. Growing student reliance on ChatGPT can promote cognitive offloading, undermining authentic learning and impeding the development of critical thinking skills (Kosmyrna et al., 2025; Zhou, Pankiewicz et al., 2025). Students often “game” agents by seeking direct answers rather than developing learning strategies. In many situations, they prioritize task completion and performance metrics (e.g., task scores) over conceptual understanding and the development of metacognitive processes (Baker, Corbett, Koedinger and Wagner, 2004; Snyder, 2024; Snyder, Hutchins, Cohn, Fonteles and Biswas, 2024). Such behaviors can create the appearance of mastery without real conceptual understanding, diverging from social-cognitive principles that emphasize sense-making and self-regulation (Bandura, 2001).

Many contemporary LLM agents emphasize knowledge acquisition but often overlook social-cognitive factors that highlight the interplay between cognition, behavior, and the environment. For example, SCT identifies self-efficacy as a key driver of motivation and perseverance: learners become more autonomous when they attribute their progress to their own efforts and strategies rather than relying on assistance (Schunk and Usher, 2012; Schunk, 1994; Ponton and Rhea, 2006). Therefore, pedagogical agents must enhance competence and confidence without encouraging dependence.

✉ clayton.a.cohn@vanderbilt.edu (C. Cohn)
ORCID(s): 0000-0003-0856-9587 (C. Cohn)

This requires interactions be *personalized* to students and *adaptive* over time. Relying solely on student-agent chat histories for agent decision-making prevents this, as agents lack awareness of students' actions in the environment. Multimodal learning analytics (MMLA) emphasizes contextualizing student-agent interactions with trace-log evidence, enabling inferences not possible from chat histories alone (Ouyang and Bai, 2026; Zhang, Borchers, Cohn, Srivastava, Snyder, Guo, TS, Mohammed, Noh and Biswas, 2026; Cohn, Rayala, Snyder, Fonteles, Jain, Mohammed, Timalsina, Burriss, T S, Srivastava, Deweese, Eeds and Biswas, 2025b).

Recently, the AI community has adopted *agentic* AI, in which agents autonomously plan, reason, and execute tasks to achieve defined goals, often within *multi-agent* architectures (Stryker, 2025). These systems reason over multimodal data, modularize components for specialized tasks, and provide interpretability by tracing calls between modules. Such transparency fosters trust among students and teachers, who are often reluctant to rely on opaque pedagogical systems (Khosravi, Shum, Chen, Conati, Tsai, Kay, Knight, Martinez-Maldonado, Sadiq and Gašević, 2022; Tsai, Whitelock-Wainwright and Gašević, 2021; Hakami and Hernández Leo, 2020; Cohn, T S, Mohammed and Biswas, 2025d).

Understanding how students interact with LLM agents is essential for aligning agent feedback with learning theory, teacher intent, and curricular goals. This calls for authentic, in situ classroom studies that reveal patterns of how students interact with and engage these agents over time. Such evidence is necessary to design pedagogical agents that move beyond rudimentary task completion and performance optimization to cultivate higher-order, transferable problem-solving skills (Wu, Zhai and Song, 2025; Ganguly, Mehjabin, Malik and Johri, 2026).

To address these needs, we developed an agentic, multi-agent, multimodal Collaborative Peer Agent, *Copa*, based on the Evidence-Decision-Feedback (EDF) framework (Cohn, Guo, Rayala, Wang, Mohammed, Timalsina, Jain, Eeds, Deweese, Osborn Popp, Stanton, Walker, Ma and Biswas, 2026; see Section 3). *Copa* combines capabilities that were previously fragmented across different pedagogical agents by incorporating theory-guided scaffolding, multimodal learner modeling derived from environmental logs and chat histories, multi-agent reasoning for modular decision-making, and traceable feedback generation to enhance interpretability. We deployed *Copa* in a high school science classroom with $n = 33$ student dyads learning physics and computing through computational modeling, where *Copa* facilitated exploration, experimentation, and problem-solving.

We evaluated whether *Copa* adapts its scaffolding, supports conceptual articulation, avoids fostering dependence, builds confidence, remains interpretable, and yields distinct interaction trajectories over time. Overall, findings were positive: *Copa* adapted its support as learners progressed, students increasingly verbalized conceptual understanding and confidence, reliance on the agent decreased with mastery, the system's responses remained traceable to multimodal evidence, and distinct student-agent interaction patterns emerged. Taken together, this paper contributes both technically and pedagogically by introducing a theory-guided, agentic, multi-agent, multimodal LLM agent, and by showing in an authentic classroom setting how such an agent can support adaptive STEM+C learning.

2. Background

Prior to LLMs, ITS frameworks such as *Adaptive Control of Thought-Rational* (ACT-R) (Ritter, Tehranchi and Oury, 2019) and *Knowledge-Learning-Instruction* (KLI) (Koedinger, Corbett and Perfetti, 2012) modeled human cognition to predict task performance and align instruction with cognitive processes, enabling precise model-tracing feedback. ITS systems were intentionally narrow in scope, prioritizing depth over breadth to methodically guide students towards domain mastery and systematic learning. Accordingly, agents in these environments operated within discrete state and action spaces, relying on predefined models of student knowledge, pedagogical strategies, and domain content to support student learning.

More recently, open-ended learning environments (OELEs) enable learners to construct their own understanding through exploration, experimentation, and problem-solving via paradigms such as “learning by teaching” (Munshi, Biswas, Baker, Ocumpaugh, Hutt and Paquette, 2023) and “learning by modeling” (Zhang, Biswas, McElhaney, Basu, McBride and Chiu, 2020). Agents in these environments are well-suited for LLMs, as they operate in continuous spaces and support open-ended tasks. ACT-R and KLI are less well-suited to such contexts, as they are not designed to address the complexities of open-ended environments and the opacity of LLM systems.

Since the advent of LLM pedagogical agents, little attention has been paid to frameworks that are both grounded in learning theory and well-suited to LLM adoption. Simultaneously, the ease of developing LLM agents (e.g., via “vibe coding” and API calls) has driven a proliferation of systems leveraging advances such as multi-agent architectures, multimodal learner modeling, and agentic retrieval-augmented generation (RAG) (Hou, Forsyth, Andrews-Todd,

Paper	Theoretical Foundation	Agentic	Multi-Agent	Multi-modality	Authentic Classroom Study	Personalization	Adaptability	Interpretability	Interaction Patterns	Reliance	Self-Efficacy	Temporality
Aulia et al. (2025)	✓	✗	✗	✓	✗	✓	✓	✗	✗	✓	✗	✗
Cohn et al. (2025c)	✓	✗	✗	✗	✓	✓	✓	✓	✗	✗	✓	✓
Chu et al. (2025)	✗	✗	✓	✗	✗	✗	✗	✓	✗	✗	✗	✗
Dai et al. (2025)	✗	✓	✓	✗	✗	✓	✗	✗	✗	✗	✗	✗
Hou et al. (2025)	✓	✗	✓	✗	✓	✗	✗	✓	✓	✗	✗	✗
Jin et al. (2025)	✓	✗	✗	✓	✓	✗	✗	✗	✗	✗	✗	✓
Li et al. (2024)	✗	✗	✓	✗	✗	✓	✓	✗	✗	✗	✗	✗
Liu et al. (2025b)	✓	✗	✗	✗	✓	✗	✓	✗	✗	✗	✓	✗
Liu et al. (2025a)	✓	✗	✓	✗	✗	✓	✓	✗	✗	✗	✗	✗
Scholz et al. (2025)	✗	✓	✓	✗	✗	✓	✓	✗	✗	✓	✓	✗
Shi et al. (2025)	✓	✗	✓	✗	✗	✗	✓	✗	✓	✗	✗	✗
Sixu et al. (2024)	✗	✗	✗	✓	✓	✗	✓	✗	✗	✗	✗	✗
Sun et al. (2025)	✗	✗	✓	✗	✗	✗	✗	✗	✗	✗	✗	✗
Wang et al. (2025)	✗	✗	✓	✓	✓	✓	✓	✗	✓	✗	✗	✗
Xi et al. (2025)	✓	✗	✗	✗	✓	✓	✓	✗	✓	✗	✗	✓
Zha et al. (2025)	✗	✗	✗	✓	✗	✓	✓	✓	✓	✗	✗	✗
Zhang et al. (2025)	✗	✗	✓	✗	✗	✓	✗	✗	✗	✗	✗	✗
Ours (2026)	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓

Table 1

LLM pedagogical agents in prior research, mapped across design features, capabilities, and study foci.

Legend. TRUE=✓, FALSE=✗.

Theo. foundation: explicit grounding in learning theory.

Multimodality: includes non-textual inputs (e.g., logs).

Adaptability: scaffolding evolves with learners.

Reliance: considers student reliance on agent.

Agentic: explicitly agentic.

Auth. class. study: in-situ study w/ classroom students.

Interpretability: feedback traceable to input data.

Self-efficacy: considers student confidence or self-efficacy.

Multi-agent: explicit multi-agent architecture.

Personalization: agent offers tailored interactions.

Interact. patt.: analyzes student-agent interaction patterns.

Temporality: longitudinal evaluation (over time).

Rice, Cai, Jiang, Zapata-Rivera and Graesser, 2025; Di Mitri, Srivastava, Fernandez Nieto, Echeverria, Cobos, Tudur Sadashiva, Spikol, Wong, Zhou and Cukurova, 2025; Vatal, Biswas and Goldberg, 2023; Cohn et al., 2026).

Recent work reflects a trend toward prioritizing technical novelty over pedagogical benefit, as several agents are “theoretical” architectures without classroom evaluation (Scholz, Nguyen, Singla and Nagashima, 2025; Li et al., 2024; Dai, Jiang, Chen, Guo, Liu and Shao, 2025), limiting insight into how students interact with these systems or whether they improve learning over time. Even among systems deployed in classrooms, most rely solely on chat histories rather than multimodal data, constraining feedback personalization, adaptive scaffolding, and holistic learner modeling (Chu, Li, Yang, Shomer, Liu, Copur-Gencturk and Tang, 2025; Shi, Liang and Xu, 2025; Cohn, Rayala, Srivastava, Fonteles, Jain, Luo, Mereddy, Mohammed and Biswas, 2025c). Despite recent technical gains, the current landscape of LLM pedagogical agents remains underdeveloped in several key areas.

As summarized in Table 1, prior LLM pedagogical agents typically address only subsets of the dimensions we examine (Sun and Tai, 2025; Li et al., 2024; Liu et al., 2025b). Agents may be technologically advanced but lack theoretical grounding or attention to student experience (Chu et al., 2025); they may be pedagogically principled but lack interpretability (Aulia et al., 2025); and studies may consider agent adaptation and personalization but without addressing social-cognitive factors such as self-efficacy (Khosrawi-Rad, Keller, Benner, Grogorkick, Borchers, Janson, Leimeister and Robra-Bissantz, 2025) and agent reliance (Li et al., 2024). Here, we address these dimensions jointly.

3. Theoretical Framework

In response to limitations of pre-LLM, ITS-based frameworks, Cohn et al. (2026) proposed a framework for pedagogical agent scaffolding in the LLM era that delivers feedback (1) grounded in learning theory, (2) interpretable to stakeholders, and (3) personalized and adaptive to students.

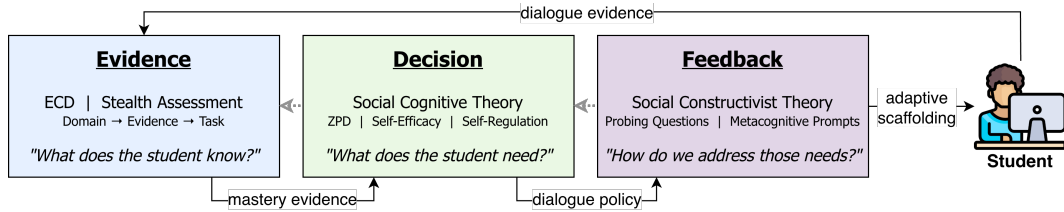


Figure 1: Evidence-Decision-Feedback (EDF) framework (Cohn et al., 2026).

As shown in Figure 1, the **Evidence-Decision-Feedback (EDF)** framework comprises three semi-autonomous modules that jointly operationalize adaptive student support. Grounded in evidence-centered design (ECD) (Mislevy, Almond and Lukas, 2003), the **Evidence** module specifies the student data to be monitored and updates the learner model accordingly. In practice, this process draws on Stealth Assessment (Shute, 2011), which unobtrusively assesses students' knowledge gain and development of cognitive and metacognitive processes as they progress through their learning during problem-solving tasks, enabling the agent to define and analyze target indicators without disrupting students' learning activities.

The **Decision** module consumes *mastery evidence* produced by the Evidence module to determine pedagogical intent (such as strengthening student confidence in kinematics and computational modeling) and generate an SCT-aligned *dialogue policy* within the student's ZPD. The resulting policy specifies the type and level of support needed, calibrating scaffolding to meet students' needs while encouraging them and promoting self-regulation. The **Feedback** module implements this policy by delivering adaptive scaffolding grounded in social constructivist principles (e.g., support for active learning and reflection), translating pedagogical intent into concrete conversational strategies consistent with the selected dialogue policy.

Gray dotted arrows in Figure 1 denote *interpretability*, a core tenet of EDF, defined by Cohn et al. (2026) as “the extent to which agent output can be traced through the input-to-evidence-to-policy-to-feedback chain using observable system artifacts, thereby supporting stakeholders’ ability to understand why the agent produced a given output.” In LLM-based systems, interpretability can be achieved by requiring each module to use chain-of-thought reasoning and justify its decisions based on the inputs it receives.

4. Collaborative Peer Agent (Copa)

Copa is an agentic, multi-agent, multimodal **Collaborative Peer Agent** for STEM+C learning powered by GPT-5. It adopts a *knowledgeable peer* persona, preferred by students who value the emotional support of a peer and the competence of an expert (Cohn, Fonteles, Snyder, Srivastava, T S, Campbell, Montenegro and Biswas, 2025a). This role promotes inquiry and exploration while aiming to reduce the dependence that can occur with expert tutors (Moos and Azevedo, 2009).

Copa operates agentially, coordinating sub-agents to gather evidence, identify learner needs, determine students' ZPD, retrieve domain knowledge, and deliver scaffolding within the EDF framework. Copa reasons over student log and chat data, selects actions, and generates feedback under defined constraints. This reflects *bounded autonomy* (Pan, Arabzadeh et al., 2025), which balances agent initiative with pedagogical guardrails (e.g., *human-in-the-loop* methods (Memarian and Doleck, 2024; Alfredo, Echeverria, Jin, Yan, Swiecki, Gašević and Martinez-Maldonado, 2024; Office of Educational Technology, 2023; Fonteles, Cohn, Ayalon, Zhou, T.S., Davalos, Li, Rayala, Mereddy, Coursey, Jain, Zhang, Enyedy, Danish and Biswas, 2026; Cohn, Snyder, Montenegro and Biswas, 2024c)) to support learning while preserving trust and interpretability. Rather than relying on unconstrained long-horizon planning, agent behavior is structured around the three EDF modules (discussed shortly).

Copa is embedded within *C2STEM* (Hutchins, Biswas, Maróti, Lédeczi, Grover, Wolf, Blair, Chin, Conlin, Basu et al., 2020) (see Figure 2), a complex STEM+C OELE targeting one- and two-dimensional kinematics. Students work in dyads to build *computational models* that simulate the motion of various objects (e.g., a truck or drone) using kinematic equations and block-based code, and interact with Copa at their discretion. Tasks are open-ended, with multiple valid solution paths and no single correct answer.

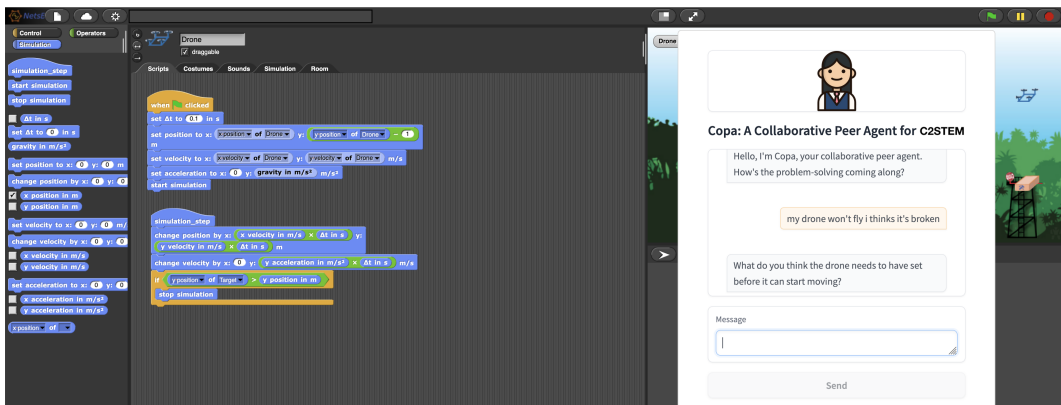


Figure 2: Copa, pictured within the C2STEM learning environment.

C2STEM interactions generate log data to support agent reasoning, including *logged actions* (raw sequences such as adding, editing, or removing blocks), *model state* (the code blocks currently on screen), and *task context* (the component under development, e.g., variable initialization). Logged actions are transformed into *processed actions* to provide semantically meaningful representations to the LLM (e.g., `set_block1234_vel_4` $\rightarrow$ Set Velocity to 4 m/s). Logs also support progress tracking by defining rubric criteria (e.g., correctly setting the initial velocity to 0) for each problem-solving task, then translating students’ models into a canonical form that enables evaluation.

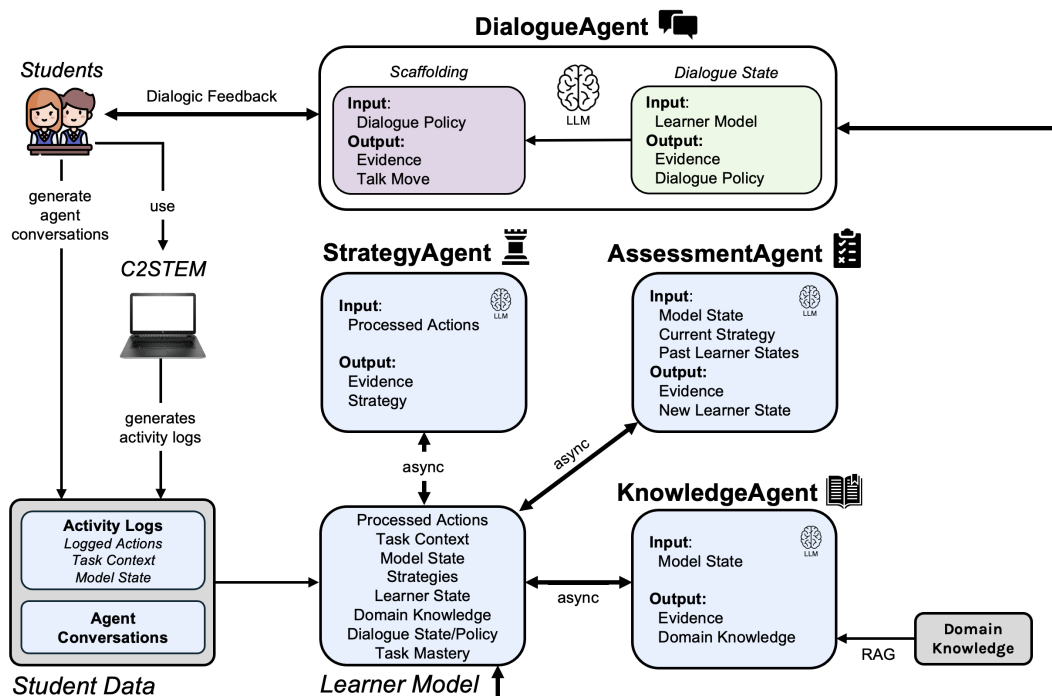


Figure 3: Copa and its four sub-agents. Colors correspond to the EDF modules presented in Figure 1. Gray boxes = data stores. Thick edges = two-way relationships.

Shown in Figure 3, Copa comprises four specialized sub-agents operating over a shared *learner model* that is continuously updated through *activity logs* and *agent conversations*. The learner model includes features such as students’ actions, model state, and task mastery (i.e., a running model score), which are illustrated in Figure 3 and described throughout this subsection. Copa’s sub-agents map directly onto the EDF modules through color-coding—with emphasis on the *Evidence* module (blue in Figure 3), which comprises three asynchronous agents—the StrategyAgent, AssessmentAgent, and KnowledgeAgent—that transform noisy log data into a coherent learner model.

The **StrategyAgent** uses a sliding window approach to reason over the students' last ten C2STEM environment actions (e.g., ...BUILD → ADJUST → EXECUTE → ADJUST) to identify the current *strategy* they are employing in their model-building and debugging tasks. StrategyAgent distinguishes productive strategies (e.g., TINKERING: systematic trial-and-error) from unproductive approaches (e.g., DEPTH-FIRST ENACTING: building without periodic testing) to gain deeper insight into students' cognitive processes. Prior work on self-regulated learning (SRL) emphasizes representing actions temporally rather than as isolated events (Snyder, 2024). These strategies capture *how* students approach problem-solving, not only *what* they do on screen. This provides structured, higher-level representations for decision-making that transcend the noise of individual actions.

The **AssessmentAgent** updates the *learner state* (e.g., DEBUGGING, PROGRESSING) by reasoning over the current model state and recent strategy identified by the StrategyAgent. For example, if the model has no errors but the strategy used to construct recent segments is suboptimal, students are categorized as EXPLORING; if they cannot resolve a bug after multiple actions and interactions, they are categorized as STRUGGLING. Consistent with Stealth Assessment, the AssessmentAgent continuously and unobtrusively evaluates student performance within the environment. By integrating the on-screen computational model with inferred strategies, AssessmentAgent provides an evidence-based account of students' progress in the current problem-solving phase, ensuring agent decision-making and feedback are aligned with their immediate needs.

The **KnowledgeAgent** compares the current model state to an expert reference (i.e., a Python-like textual representation of a full-credit model). Using the difference between students' current computational models and the expert reference, KnowledgeAgent identifies conceptual knowledge gaps to approximate their ZPD and retrieve relevant *domain knowledge* (i.e., kinematics, computing, and C2STEM content) from a knowledge base. The knowledge base is a vector store of curriculum-aligned information derived from prior empirical accounts of student difficulties in C2STEM. Retrieval is performed using an agentic RAG approach, LC-RAG (Cohn et al., 2025b). Rather than relying on semantic alignment between the student query and the vector store, the KnowledgeAgent reasons over the gap between the students' current and expert models and specifies which information to retrieve. This creates a clear semantic mapping to the knowledge base that is often absent from students' queries because they may struggle to understand or articulate their needs. For example, if students fail to update velocity at each simulation step, the agent recognizes this and retrieves the appropriate kinematic equations, along with guidance on variable updates and loops, and stores this information in the learner model to support subsequent tailored feedback.

The **DialogueAgent** functions as the central orchestrator, integrating the *Decision* (green) and *Feedback* (purple) modules. Because these modules operate synchronously¹, consolidating them within a single agent reduces latency while allowing Evidence agents to perform asynchronous reasoning without interrupting conversational flow. Upon receiving a student query, the DialogueAgent infers the student's *dialogue state*, a discrete representation of the student's current conversational status (e.g., DEMONSTRATES_UNDERSTANDING, where students correctly verbalize conceptual knowledge to Copa; and INFORMATION_SEEKING, where students request information).

Once the dialogue state is determined, the DialogueAgent integrates it with the learner model to determine pedagogical intent and generate a *dialogue policy*, i.e., the agent's optimal next action, such as PROBE_UNDERSTANDING (when students request direct solutions or express misconceptions, prompting Copa to probe their kinematics and computing knowledge); SUGGEST_ACTION (prompting a concrete action in C2STEM as a hint toward advancing their model); or PUSH_LIMIT (encouraging students to extend their understanding beyond their demonstrated task mastery). The DialogueAgent operationalizes this policy by producing a *talk move* (a response to the students) and is instructed to “*promote self-efficacy through encouragement and praise*” and ask “*probing questions*,” consistent with social-cognitive and social-constructivist principles.

All agents provide *evidence* with their outputs, using chain-of-thought reasoning to justify decisions and support interpretability pursuant to EDF. For example, when students asked, “*is our start okay?*” the DialogueAgent reasoned over multimodal inputs and learner model information from the Evidence agents, generating the following rationale (i.e., evidence), dialogue policy, and response:

Agent Rationale: The students are checking if their initialization setup is correct...it's best to probe their understanding of why the initial position value matters before confirming correctness...

Selected Dialogue Policy: PROBE_UNDERSTANDING

Response to Students: How does the starting position value affect where the truck begins its motion?

¹This is required because the agent cannot know what students will say in advance. Copa needs a student query to generate a dialogue state, which the DialogueAgent uses to derive a dialogue policy. That policy is necessary for the formulation and delivery of feedback.

All C2STEM task rubrics, our model selection procedure, and prompt engineering procedures are detailed in our [Supplementary Materials](#)², as are all agents' prompts, inputs, outputs, functionalities, and theoretical couplings.

5. Methods

To evaluate Copa, we developed the following research questions (RQs):

- RQ1. How does Copa's scaffolding strategy (dialogue policy selection) change as students' task mastery increases?
- RQ2. How does students' verbally-demonstrated conceptual understanding during Copa interactions change as they progress in their problem-solving?
- RQ3. How does students' reliance on Copa change as task mastery increases?
- RQ4. How does students' expressed confidence in agent dialogue change with task mastery during computational modeling?
- RQ5. To what extent is Copa's evidence-decision-feedback chain grounded in student interaction data (logs and dialogue)?
- RQ6. What contrasting student-agent interaction trajectories can be observed among dyads across computational modeling sessions, and how are these trajectories associated with problem-solving strategies and learning outcomes?

These RQs assess Copa through three different lenses:

1. *Agent behavior* during problem solving (RQ1, RQ5),
2. *Student behavior* and engagement while interacting with Copa (RQ2, RQ6),
3. *Learner motivational factors*, including student reliance on Copa and verbalized confidence (RQ3, RQ4).

Each corresponds to a principle derived from the framework described in Section 3, and its theoretical motivation is explained at the beginning of the relevant subsection to avoid repetition and repeated references to the RQs.

In collaboration with five STEM educators and cognitive scientists (hereafter, "the Educators"), we deployed Copa in an authentic high school science classroom in the southeastern United States with $n = 33$ sophomore dyads³ (ages 15-16; 48% male, 48% female, 2% no response; 68% White, 28% Hispanic, 20% Black, 12% Asian, 4% Native American)⁴ participating in the C2STEM kinematics curriculum over six weeks. Sessions met weekly for 2 hours and were taught by this paper's lead author with educator oversight.

During the study, students completed two warm-up tasks to build familiarity with the agent. Analyses focus on three subsequent tasks following this acclimation period: (1) a 1-D (*linear motion*) *Truck Task* where students model a truck speeding up, cruising, slowing down, and stopping; (2) a 2-D (*parabolic motion*) *Drone Task*, where students program a drone model to drop a package on a target; and (3) a 2-D *Two-Package Drone Task* where students program a drone model to drop two packages on two different targets. Each session began with instruction on one- and two-dimensional kinematics, connecting physics concepts (e.g., position, velocity, acceleration, Δt) to computational constructs (e.g., variable initialization and updating, loops, conditionals), followed by modeling activities in C2STEM.

We collected 7,017 logged environment actions (mean ≈ 213 per dyad per session) and 238 student-agent conversation turns (mean ≈ 7 per dyad per session) across the three tasks, along with the CoT reasoning (i.e., "evidence" in Figure 3) associated with each agent decision. Screen recordings, student video, and dyadic conversations were captured using webcams and lapel microphones via an internal open-source multimodal data collection and alignment tool, SyncFlow (Timalsina, Davalos_Anaya, Sanda, Zhang, Horn_Fonteles, T_S and Biswas, 2025). Pre- and post-tests were administered at the study's outset and conclusion, respectively. At the end of the study, students completed an anonymous semi-structured exit survey, and all findings were discussed with three of the participating Educators, who provided additional feedback. Together, these data streams constitute the dataset analyzed in this paper.

²https://github.com/claytoncohn/CE26_Supplementary_Materials

³All analyses were conducted at the dyad level (one pair of students per computer).

⁴Percentages exceed 100% because some students identified as multiple ethnicities.

Dialogue Policy	Spearman's ρ	Trend	p -value
PROBE_UNDERSTANDING	-0.34	Decreasing	0.034
SUGGEST_ACTION	0.33	Increasing	0.039
PUSH_LIMIT	0.42	Increasing	0.007

Table 2

Copa's dialogue policy (see Section 4) adaptation across task mastery quintiles.

To address the research questions, we adopted a mixed-methods approach integrating quantitative and qualitative analyses (detailed in the following section). Quantitative analyses were conducted computationally and are reported using predefined metrics. Because the sample size and classroom context did not support a randomized controlled trial (RCT) with experimental and control conditions, the study is exploratory; accordingly, reported quantitative associations are correlational rather than causal. All authors, who participated in conducting the study and data collection, were involved in the qualitative analyses of the results. We manually reviewed screen recordings, student videos, collaborative discussions, agent-interaction transcripts, and log data. We used memoing (Hatch, 2002; Birks, Chapman and Francis, 2008) to document key findings and to support ongoing team-based interpretations. We triangulated all findings across data modalities (e.g., aligning agent dialogue with peer discourse and on-screen activity) to substantiate inferences. All study participants and their parents provided informed assent/consent, and all procedures were approved by Vanderbilt University's IRB.

6. Findings

In this section, we describe the experimental design for each RQ and present the quantitative and qualitative findings. We distinguish between students' *understanding*, i.e., their verbalized knowledge during conversations with the agent that demonstrates what they know, versus their *task mastery*, i.e., their performance measured as the percentage rubric score achieved on a given C2STEM task.

6.1. RQ1: How does Copa's scaffolding strategy (dialogue policy selection) change as students' task mastery increases?

Consistent with the Zone of Proximal Development, instructional scaffolding should be dynamically calibrated and progressively faded as learners internalize skills and achieve independent performance. Initially, Copa should elicit and diagnose students' understanding to estimate their current domain knowledge and ZPD by systematically identifying knowledge gaps and misconceptions. As task mastery develops, Copa should transition to providing increasingly specific conceptual and procedural guidance that supports learners in extending their reasoning beyond their current level of competence.

To evaluate Copa's adaptivity (RQ1), we examined how its dialogue policies evolved as students progressed in their model-building tasks by analyzing the frequency of policy selection across quintiles of students' task mastery (see Table 2). Quintiles enabled stable comparisons of agent behavior across meaningful phases of learning. Policy frequencies were normalized to the percentage of total dialogue each dyad was involved in and averaged across quintiles. For each dialogue policy, we computed Spearman's ρ (given the ordinal mastery data) to assess the association between shifts in dialogue policy and task mastery.

Table 2 shows that Copa significantly⁵ reduced its use of the PROBE_UNDERSTANDING policy as mastery increased ($\rho = -0.34$, $p = 0.034$), while SUGGEST_ACTION ($\rho = 0.33$, $p = 0.039$) and PUSH_LIMIT ($\rho = 0.42$, $p = 0.007$) increased in relative frequency.

Consider the following example in the one-dimensional C2STEM Truck Task. One dyad reported, "*The truck isn't starting at -60 meters*" at 0% mastery. Copa responded with a PROBE_UNDERSTANDING utterance: "*What happens if the starting values are set after the simulation has already begun running?*," which targeted learners' understanding of initialization rather than directly correcting the error. Later, at 65% mastery, the students stated, "*We have successfully developed a code that stops the truck at the stop sign.*" Copa shifted to the SUGGEST_ACTION policy, i.e., "*Want to try using the speed limit variable instead of the number for cruising?*" to prompt use of the SPEED_LIMIT constant instead of hard-coding a value. After task completion, Copa escalated to the PUSH_LIMIT policy: "*Your truck model works perfectly! Now let's tweak some values and see what happens!*," encouraging further exploration.

⁵Statistical significance is assessed at $\alpha = 0.05$ throughout the paper.

This example, in line with ZPD principles, illustrates how a knowledgeable peer (Copa) facilitates the transition from eliciting and assessing conceptual understanding to supporting its application by providing hints to support model construction, and then directing assistance toward challenges beyond learners' independently demonstrated competence. Overall, these findings show that **Copa's scaffolding strategy adapted to increasing task mastery by shifting from assessing students' knowledge early on to providing more concrete guidance and extension-oriented support as they progressed.**

6.2. RQ2: How does students' verbally-demonstrated conceptual understanding during Copa interactions change as they progress in their problem-solving?

Students often exhibit gaming behaviors in computer-based learning environments, completing tasks without understanding the targeted concepts or their relationships (Cock, Marras, Giang and Käser, 2022). Therefore, task mastery alone may overestimate learning in our setting, as students can progress through trial and error without developing a principled understanding of the domain concepts and relations governing their computational model code for object motion. Social Constructivism posits that learning is socially mediated and evidenced not only by performance but also by learners' ability to articulate, justify, and negotiate meaning through dialogue. Accordingly, we examine the relationship between task mastery and verbalized understanding during Copa interactions (RQ2) to assess whether observed progress reflects substantive learning.

We analyzed students' response *success rates*, defined as the proportion of instances in which students correctly explained their code to Copa via a DEMONSTRATES_UNDERSTANDING dialogue state immediately following a PROBE_UNDERSTANDING policy, as a function of task mastery deciles. Deciles enabled a finer-grained assessment of how verbally-demonstrated understanding varied across incremental mastery levels in C2STEM. To mitigate bias arising from interaction frequency, data were normalized by computing, for each dyad, the mean response success rate across mastery deciles and then correlated using Spearman's ρ .

Figure 4 shows that the success rate increased with mastery ($\rho = 0.40$, $p = 0.014$).⁶ This association suggests that **students' in-environment progress co-occurred with an improved ability to verbalize domain concepts and their relationships to Copa**—consistent with gains in conceptual understanding rather than superficial task optimization, in line with social constructivist principles.

For example, at the start of the Truck Task (i.e., 0% demonstrated mastery), one dyad asked “*Is our script correct so far?*” in the hope that Copa would identify any errors in their model. Instead of providing a direct answer, Copa prompted the dyad to consider missing initializations: “*What other quantities do we need to set before the simulation can start correctly?*” using a PROBE_UNDERSTANDING dialogue policy. The students deflected rather than answering the question, again attempting to elicit an answer: “*I meant our delta t, like is it supposed to be .1*”. Later in the same session, when the dyad had attained 45% mastery, Copa asked “*What quantities do you think connect how fast the truck is going to how far away it needs to start slowing down?*” This time, the students replied “*acceleration and velocity*,” correctly identifying the relevant physics quantities, which Copa classified as DEMONSTRATES_UNDERSTANDING. This shift from deflection to articulation of domain relationships illustrates how gains in task mastery coincided with improved verbalization of the underlying physics concepts.

6.3. RQ3: How does students' reliance on Copa change as task mastery increases?

During participatory design, students reported wanting control over when and how they engaged with Copa, highlighting the importance of *learner agency* during agent interactions (Cohn et al., 2025a). However, increased control risks promoting over-reliance if students engage in cognitive offloading rather than developing the capacity

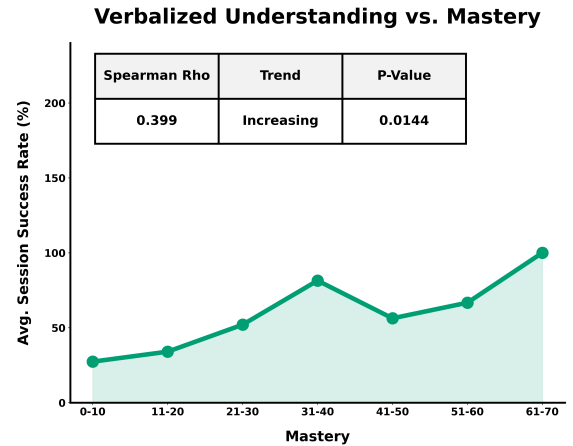


Figure 4: Correlation between success rate and mastery deciles.

⁶Success rate for mastery > 0.7 was not computed due to insufficient PROBE_UNDERSTANDING policies: once mastery was high, Copa reduced probing because students no longer requested direct answers or exhibited misconceptions.

to learn and apply shared knowledge. As students progress in their problem-solving, they should be relying less on Copa and more on their shared knowledge. Consistent with SCT, such a shift reflects growth in self-regulation and self-efficacy in applying domain knowledge, and motivates evaluating whether Copa functions as productive support or a persistent crutch during task progression.

To evaluate RQ3, we examined how the frequency of students' engagement with the agent changed across task mastery deciles. We computed the percentage of agent support (i.e., interactions) within each decile to quantify the proportion of Copa interactions per session. To control for imbalanced interaction volumes, we normalized the data by averaging the agent-support proportion for each dyad across deciles, ensuring high-frequency dyads did not exert disproportionate influence. We again used Spearman's ρ to correlate agent interaction frequency with task mastery.

Figure 5 shows that requests for agent support were high early in learning: 59% occurred when mastery was low (< 40%). **As mastery increased, Spearman's $\rho = -0.26$ ($p < 0.001$) showed moderate decreases in requests for agent support, indicating reduced reliance on the agent over time.** Triangulation of student discourse, screen recordings, and video data further showed that, as mastery increased, students drew more on their own knowledge and partner collaboration, discussion, and debate than on agent support.

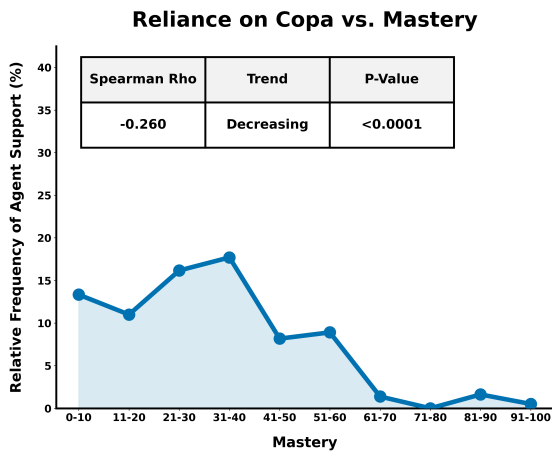


Figure 5: Correlation between student requests for agent support and mastery deciles.

rate of descent and velocity. Is our code good now?," at which point Copa realized that their "model still has issues with how PACKAGE and PACKAGE2 update velocity and position," asking students "How should the packages' velocity change once gravity starts acting on them?"

Several other dyads displayed similar behavior. Students relied heavily on Copa when they began constructing their models, but increasingly discussed the agent's probing questions with their partner or reflected independently, drawing on their own reasoning rather than requesting further help. In many cases, C2STEM progress occurred simply by translating Copa's prompts into model construction and debugging actions through peer discussion. This suggests that by asking probing questions rather than providing direct answers, Copa prompted students to invoke their own reasoning about problems that they might otherwise have delegated to the agent. This is consistent with SCT's emphasis on the gradual internalization of support and growing self-regulation. We discuss broader implications for the optimal "helpfulness" of pedagogical agents in Section 7.

6.4. RQ4: How does students' expressed confidence in agent dialogue change with task mastery during computational modeling?

SCT highlights the importance of self-efficacy and its relationship to students' engagement, persistence, and academic performance. In this study, it is essential to examine how confident or uncertain students are in their computational modeling, as this directly influences their problem-solving abilities. As students engage in the C2STEM

We illustrate this pattern with a dyad working on C2STEM's Two-Package Drone Task. At the outset (0% mastery), one student asked Copa, "how should I calculate the y displacement equation?," seeking a direct answer rather than trying to reason through the problem by himself. Copa responded by probing the dyad's understanding: "What quantities do you think affect how far something moves vertically over time?," noting that, because the student was just beginning and unsure how to express vertical displacement, it was preferable to elicit relevant factors governing vertical motion before introducing equations to represent the 2-D motion.

Rather than replying to the agent, the student paused and instead voiced his reasoning to his partner: "I think...gravity is how far something moves vertically over time." Although incorrect, the statement initiated discussion, and the pair decided "let's start with the drone." They then performed 288 consecutive independent actions: building, adjusting, and testing their model while using C2STEM's graph tool to visualize how variables changed over time, without Copa support. They reengaged the agent only for validation once their model was nearly complete: "We have made the packages drop at the right

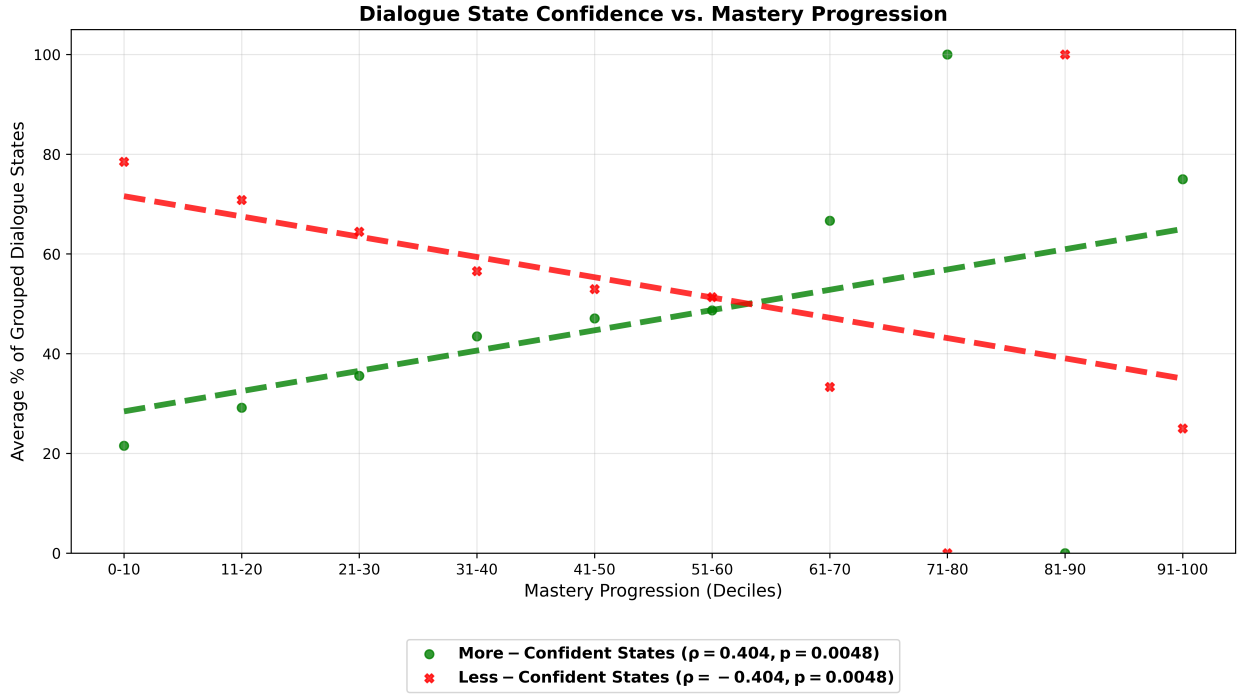


Figure 6: Student dialogue state confidence across mastery deciles.

tasks and interact with Copa, they should become more confident in their ability to build, test, and refine these models as their problem solving improves.

To evaluate RQ4, we used dialogue state groupings—i.e., *more-confident* versus *less-confident*—as a proxy for confidence, correlating it with mastery. Specifically, we examined shifts from **less-confident** states (such as seeking information, expressing frustration, or not understanding Copa’s feedback) to **more-confident** states (such as suggesting concrete actions, applying the feedback provided to advance their model development in C2STEM, or correctly expressing conceptual understanding) as student task mastery improved. To eliminate interaction frequency bias, we computed these proportions for each dyad and mastery decile. We then estimated the overall association between dialogue state groupings and mastery deciles using Spearman’s ρ .

Figure 6 illustrates that as student mastery improved, their dialogue states shifted from less-confident to more-confident ($\rho = 0.40, p = 0.005$) as they grew from relying on Copa for information to articulating their own understanding and reasoning, and suggesting and implementing concrete actions in the C2STEM environment. This became even more clear when analyzed alongside video, screen, and speech data, which showed that students gained confidence not only in their interactions with Copa and C2STEM, but also their interactions with each other. Together, these findings indicate that **students became increasingly confident in their knowledge and abilities as they progressed in their computational modeling tasks.**

For example, one representative dyad (at mastery level = 14) began their Two-Package Drone session by asking Copa “hey, how does our code look so far,” a low-confidence, information-seeking utterance that outsourced model evaluation to the agent rather than forming their own assessment. Copa’s probes largely drew expressions of misunderstanding and uncertainty: “nothing?” and “we do?” As the dyad continued to build autonomously and mastery increased, their language shifted. At mastery level = 48, Copa asked what should happen to the package upon landing. The student replied “*velocity=0*,” which is a concise, correct verbalization with no hedging or request for confirmation from Copa. By the session’s end, at mastery level = 86, Copa invited the dyad to predict what would happen if the drone flew faster; the student replied “*it will overshoot*,” an independent causal inference again expressed without seeking Copa’s validation. This within-session transition, from doubting their model assessment to reasoning forward using

Trace Component	Metric	Actual	Random	p-value
Grounding (Data → Evidence)	Keyword Recall	0.43	0.21	$p < 0.001$
Alignment (Evidence → Decision)	SBERT Similarity	0.64	0.39	$p < 0.001$
Faithfulness (Decision → Feedback)	SBERT Similarity	0.48	0.24	$p < 0.001$

Table 3

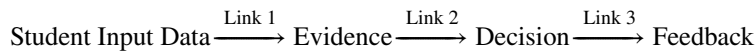
Interpretability results across all three links.

physical principles, suggests increasing confidence in their computational-modeling abilities. This interpretation aligns with SCT, which identifies mastery experiences as a key source of self-efficacy development.

6.5. RQ5: To what extent is Copa's evidence-decision-feedback chain grounded in student interaction data (logs and dialogue)?

Our previous research indicates that trust in LLM systems is influenced by their interpretability, which requires that their outputs be traceable through observable decision-making artifacts (Cohn et al., 2025d) and grounded in student input data. With LLMs, we can generate chain-of-thought explanations that help stakeholders understand the reasons behind an agent's output and determine whether to accept its decisions. This is crucial for ECD, as validity relies on linking observed student evidence to the inferences and instructional feedback derived from it.

Agent feedback must be traceable to its dialogue policy, the dialogue policy to its evidentiary inferences, and inferences to the student input data (e.g., students' computational model state, their query to Copa, environment actions, etc.). To evaluate Copa's interpretability (RQ5), we analyze its internal reasoning chains—illustrated by the agents' "evidence" outputs in Figure 3—by decomposing them into the three links between EDF modules:



Link 1, *Grounding*, measures the extent to which Copa's evidence extraction reflects student input data. We evaluate Grounding using *keyword recall*, defined as the percentage of semantically meaningful tokens derived from student log data that appear in the extracted evidence. We apply Porter stemming to account for morphological variation (e.g., accelerates $\approx$ acceleration), retaining only tokens longer than three characters with at least one alphabetic character. Significance is assessed via a permutation test ($n = 100$), comparing observed log-evidence pairs with randomly shuffled baselines.

Link 2, *Alignment*, assesses semantic coherence between Copa's Decision (dialogue policy selection) and extracted Evidence in the learner model. We compute sentence-level semantic similarity using Sentence-BERT embeddings (SBERT; all-MiniLM-L6-v2, selected for its balance between size and performance), by encoding Copa's dialogue policy and dialogue state summaries, then comparing them via cosine similarity. Significance is evaluated via a permutation test ($n = 1000$)⁷, comparing the mean similarity of true evidence-policy pairs with shuffled baselines. Link 3, *Faithfulness*, evaluates whether Copa's student-facing Feedback maintains semantic coherence with its Decision, assessed using the same procedure as Link 2.

Table 3 demonstrates that all three links display statistically significant non-random structure ($p < 0.001$), with recall and similarity scores for true pairs being nearly double those of the baseline scores. Qualitatively, we found that Copa's reasoning chains not only accurately represent the input data and evidentiary inferences, but they are also easily understandable to humans. An example from the Drone Task is presented in Table 4.

In this example, Copa used evidence inferred from multimodal input data to determine that the students had a bug in their code and to identify the cause for the "*package flying*." Prior to the students' query (asynchronously, as discussed in Section 4), KnowledgeAgent retrieved the relevant physics and modeling information, StrategyAgent identified the students as TINKERING based on their recent actions, and AssessmentAgent determined that the students were STRUGGLING with the error. Once the students engaged Copa, the DialogueAgent identified the dialogue state (INFORMATION_SEEKING) and reasoned over the evidence to select an appropriate dialogue policy (PROBE_UNDERSTANDING), inferring that the students were conflating velocity and acceleration and that a probing question was warranted. This resulted in a personalized talk move, to which the students responded with uncertainty, thereby confirming Copa's inference that they did not understand how to distinguish between acceleration and velocity.

⁷Larger n due to fewer utterances relative to environment actions in Link 1.

	Evidence	Description
<u>Student Query</u>	"why is our package flying?"	The package is initially flying straight up and not staying attached to the drone.
<i>Task Mastery</i>	"physics 27%, computing 12%, overall 21%"	Copa has access to students' current mastery level. So far, the students are still early in the task and are demonstrating more physics knowledge than computational knowledge.
<i>Recent Actions</i>	"...ADJUST → ADJUST → BUILD → EXECUTE"	Copa sees students' most recent modeling actions, showing iterative building/adjusting/testing actions.
<i>Model State</i>	"...[set velocity to x: (x_velocity of Drone) y: (gravity)]..."	Copa sees the students' incorrect computational model on screen. The initial y_velocity should be 0, not gravity.
<i>StrategyAgent</i>	TINKERING; "...several adjustments...another build before executing again—iterative build-test cycles."	StrategyAgent reasons over students' actions to output a problem-solving strategy, in this case determining they are using TINKERING—an effective, iterative approach involving repeated testing and modification.
<i>AssessmentAgent</i>	STRUGGLING "PACKAGE sets y_velocity to gravity instead of 0..."	AssessmentAgent recognizes students are struggling to identify the error responsible for the unexpected package behavior.
<i>KnowledgeAgent</i>	"...for falling packages initialize y_velocity=0, y_acceleration=gravity, then each step update position by $velocity \times \Delta t$ and velocity by $acceleration \times \Delta t$."	KnowledgeAgent retrieves the relevant physics/modeling information needed to correct the students' implementation of gravity and velocity.
<i>DialogueAgent</i>	Dialogue State: INFORMATION_SEEKING; "The student is confused about why their package is flying upward." Dialogue Policy: PROBE_UNDERSTANDING; "...they are struggling...probe their understanding of how velocity and acceleration relate before suggesting any fix."	DialogueAgent determines the students' needs for conceptual clarification and decides to probe their understanding of velocity and acceleration before offering a concrete suggestion.
<u>Agent Response</u>	"What connects how an object speeds up or slows down to the force acting on it?"	The agent guides students to recognize that the gravitational constant refers to acceleration, not velocity.
<u>Student Response</u>	"is it velocity? or acceleration?"	The students express uncertainty about the velocity-acceleration relationship, confirming Copa's inference that this distinction is unclear to them.

Table 4

Illustrative example supporting the interpretability analysis in RQ5. The first column identifies the evidentiary data source available to Copa, where items underlined in bold are student-facing, and *italicized* items represent Copa's internal reasoning (unseen by students). The **Evidence** column presents verbatim evidence drawn from system logs or student-agent dialogue that the agent can access when generating feedback. The **Description** column provides context for how this evidence informs the agent's reasoning and response.

Consistent with ECD principles, feedback remained tied to assessment, student state, and task context, showing that **Copa's decisions were interpretable, grounded in its evidentiary reasoning and student input data.**

6.6. RQ6: What contrasting student-agent interaction trajectories can be observed among dyads across computational modeling sessions, and how are these trajectories associated with problem-solving strategies and learning outcomes?

While within-session student-agent interactions are informative, they do not capture how students' knowledge and model-building abilities change over time or how their interaction patterns evolve across the curriculum. Examining how students engage with agents *longitudinally* is, therefore, essential for assessing the longer-term effects of LLM systems and using this information to improve their functionality. This perspective is particularly important from an SCT standpoint, as repeated interactions with an agent can shape students' sense of agency, self-efficacy, and capacity for self-regulated learning through ongoing feedback and performance cues.

To examine how students engaged with Copa across the curriculum, RQ6 adopts an exploratory case-study approach, focusing on two dyads with contrasting patterns of agent use, i.e., low-versus-high agent engagement. We qualitatively analyzed each dyad's environment logs, chat histories, screen and video recordings, and collaborative discourse (audio), and memoed key findings across aligned multimodal data streams (Hatch, 2002). The aim was not to generate generalizable claims, but to surface fine-grained patterns in how students learned with Copa over time.

The low-engagement dyad produced just 7 utterances to Copa across the three C2STEM tasks (3, 1, and 3, respectively). Their interactions were brief and primarily focused on *validation*, i.e., confirming progress. The high-engagement dyad, by contrast, produced 48 utterances across the same tasks (11, 22, and 15, respectively). Their exchanges with Copa were more sustained, involving multiple turns, and were inquiry-driven, using the agent not only for validation but also for *guidance*, error understanding, and uncertainty navigation. We therefore refer to the low-engagement students as *Validation Seekers* and the high-engagement students as *Guidance Seekers*. Despite their differences in engagement, both dyads exhibited similar activity levels within C2STEM, each performing approximately 900 actions.

Validation Seekers interacted with Copa sparingly, and almost always to seek external confirmation rather than to support new reasoning. Across the curriculum, their prompts primarily asked whether their code or progress “*looked right*,” positioning the agent as a validator of completed work rather than a partner in problem-solving. Notably, low agent engagement did not reflect low task engagement. The dyad remained highly collaborative as they worked in the C2STEM environment, actively questioning each other, refining ideas, and persisting even when class time ended. At the same time, every Copa response they received across the curriculum was a probe intended to determine understanding, which students largely ignored. This case suggests that low agent engagement need not signal low task engagement, but it may limit how productively students use agent feedback when confronting uncertainty.

A representative example occurred during the Drone Task, when the *Validation Seekers* asked Copa, “*How do we know where to drop the package?*” Rather than supplying a direct answer, Copa responded with a probe: “*What factors do you think decide how far the package travels before it hits the ground?*” One student immediately rejected the prompt, saying aloud, “*That’s not what I asked it at all!*” This exchange captures the dyad’s characteristic stance toward the agent (i.e., expecting direct answers), which persisted across all three tasks. However, the dyad remained energetic and highly engaged in the C2STEM environment, working together with enthusiasm, even as they declined to consider the agent’s hints.

Guidance Seekers started similarly to the *Validation Seekers* but later evolved to treat Copa as an active problem-solving partner, using it not only for validation but also to ask causal questions, test interpretations, and express frustration. Their 48 interactions elicited a much broader range of agent responses than *Validation Seekers*, including suggested actions and motivational prompts, reflecting the dyad’s more varied and expressive use of Copa. Early exchanges (Truck Task; Day 1) were primarily answer-seeking, and they did not act on Copa’s feedback immediately, but this pattern shifted over time. By Day 2 (Drone Task), they were more willing to read, discuss, and reason through Copa’s hints together, recognizing their own mistakes in the process. By Day 3 (Two-Package Drone Task), they no longer primarily asked for answers but for help getting “*in the right direction*,” and their interactions more often followed a productive cycle of asking, discussing, enacting, and analyzing.

A clear example of this shift appeared on Day 1. When the *Guidance Seekers* reached an impasse, they told Copa, “*we need help*,” prompting the agent to acknowledge their frustration and difficulty: “*It looks like this part’s tricky—maybe ask your teacher...*” With teacher support, the students were then able to state the problem more clearly to Copa (i.e., “*it should stop going faster*”), after which Copa shifted from probing to directing: “*Nice! Try updating your condition so it checks when velocity is greater than the speed limit.*” The dyad implemented the change, tested the model, and observed that the truck now operated correctly, prompting an immediate celebratory response. This experience appeared to reshape how they engaged the agent thereafter: having seen that responding to Copa’s prompts

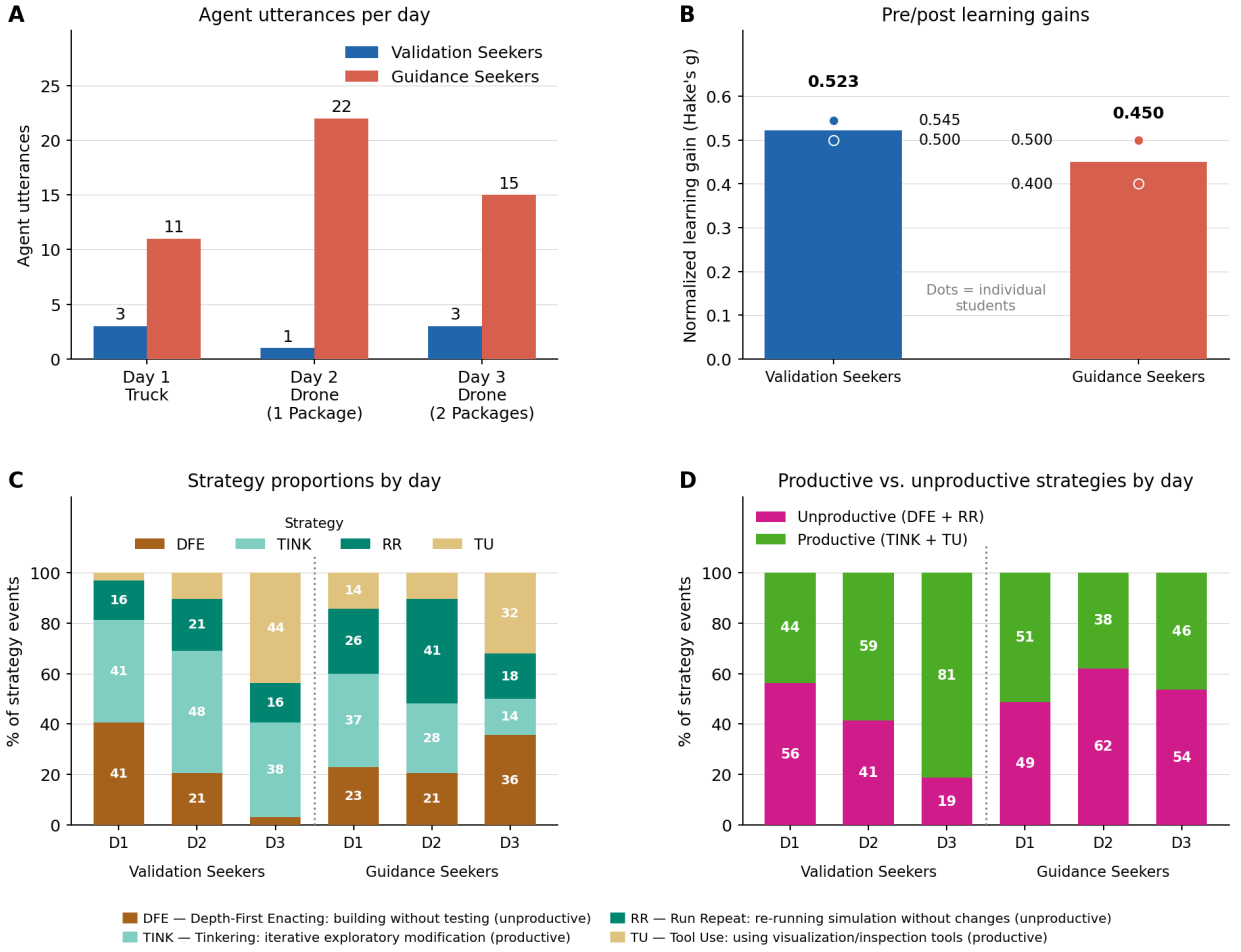


Figure 7: Comparison of the two dyads across the three sessions. (A) Agent utterances per day. (B) Pre/post normalized learning gains, with dots indicating individual students. (C) Proportions of the four problem-solving strategies by day. (D) Aggregate productive versus unproductive strategy proportions by day.

could produce tangible progress, they became more willing to treat its feedback as a resource for subsequent reasoning rather than an obstacle to a direct answer.

Figure 7 compares the two focal dyads across agent use, learning gains, and problem-solving strategies over the three sessions. As shown in Figure 7A, *Validation Seekers* interacted with Copia only rarely, whereas *Guidance Seekers* did so frequently throughout the curriculum. Despite this contrast, the dyads showed broadly similar pre/post normalized learning gains (Figure 7B), calculated using Hake gain ($g = \frac{\text{post-pre}}{\text{max-pre}}$). Their strategic trajectories, however, differed more clearly. For *Validation Seekers*, depth-first enacting declined across sessions while tool use (i.e., the graph and table tools provided in the C2STEM environment) increased, indicating a shift toward more systematic and productive problem solving (Figure 7C). This pattern is especially visible in Figure 7D, where productive strategies rose steadily from Day 1 to Day 3 as unproductive strategies fell. *Guidance Seekers* showed the opposite tendency: tinkering decreased while depth-first enacting increased across sessions, and their use of unproductive strategies tracked closely with their frequency of agent interaction. Video and audio data suggest that this pattern stemmed less from disengagement than from a rapid, build-heavy approach in which students attempted to follow the agent's guidance step by step without sufficiently testing their models between changes.

Although *Guidance Seekers* interacted with Copia far more often and in more varied ways, *Validation Seekers* showed a slightly higher mean normalized gain (0.523 vs. 0.450). Thus, richer dialogue with Copia did not translate

into stronger post-assessment performance; while sparse, validation-oriented use did not preclude substantial gains. This pattern must be interpreted in light of prior knowledge. *Validation Seekers* entered with relatively similar pretest scores (18/32 and 21/32), whereas *Guidance Seekers* were less uniform: one student began at a comparable level (20/32), but the other started substantially lower (12/32). Collaboration quality is also relevant. Although *Validation Seekers* used Copa less, the dyad worked closely together throughout the tasks, suggesting that peer interaction may have helped sustain learning even when agent use was limited. Taken together, the comparison suggests that learning gains likely reflected not only agent engagement, but also external contextual factors.

Our case-study comparison provides evidence that **student-agent interactions differed longitudinally across engagement profiles in both frequency and function, with lower-engagement students using Copa mainly for validation, and higher-engagement students using it as an active conversational partner for guidance, sense-making, and emotional support.** *Validation Seekers* collaborated well but showed limited affective and dialogic engagement with Copa, whereas *Guidance Seekers* engaged richly and emotionally with the agent but regulated strategies less consistently. Both dyads achieved similar learning gains, but through very different pathways. From an SCT perspective, this contrast suggests that student-agent interaction is shaped not only by frequency of use but by how students position the agent within the social and cognitive aspects of learning. We expound upon this in Section 7.

6.7. Student and Teacher Perceptions

Post-study survey responses indicated positive overall perceptions of Copa. Students agreed that “*the questions Copa asked were appropriate*” (mean = 3.81 on a five-point Likert scale) and that Copa “*asks questions that made me think about my model*” (mean = 3.88). These responses suggest Copa met its design goal of prioritizing critical thinking over direct solutions. For example, consider the following survey response:

When our model wasn’t working great, and we had no clue why, Copa asked specific questions that made us look deeply into the acceleration part of our code and fix the bug.

However, this behavior was not always appreciated. Ratings for perceived agent understanding (“*I feel like Copa understood what I was doing*”) and feedback utility (“*I feel like Copa’s suggestions were useful*”) were lower (mean = 2.69 and 2.77, respectively), indicating a disconnect between Copa’s pedagogical intent and students’ desire for direct answers. One student described Copa as “*broken*” and expected explicit feedback rather than ZPD-aligned hints. This reflects a broader tension between scaffolded support and student expectations, leading to frustration when direct answers are not provided.

An example of this frustration can be found in a survey response: “*One week we were very lost and asked Copa to walk us through what to do, but all COPA did was ask us questions.*” While similar survey responses were common, we also observed positive trends in the audio and video data: when students stopped seeking direct answers from Copa, they began to engage in more effective collaboration and critical thinking, often leading to the development of solutions on their own. Although frustration was evident, we view this behavior as beneficial because it facilitates self-efficacy through perseverance and negotiating meaning through discussion—while reducing over-reliance on the agent. We discuss this further in Section 7.

Findings aligned with the Educators’ classroom experiences, who reported that (1) students become frustrated when they do not receive immediate answers (“*I agree with your conclusion...students wanting immediate answers*”); (2) this expectation is shaped by prior experience with ChatGPT (“*ChatGPT [is] so desperate to provide help...they’ve already become accustomed*”); and (3) a persistent tension exists between instructional goals and student preferences, as students often “*just [want]...the right answer...check the box and move on.*” The Educators also emphasized that students hold different expectations of agents than of humans, highlighting the need for student buy-in and suggesting that sharing prior study findings with new cohorts can communicate Copa’s benefits at the study’s outset.

7. Discussion and Conclusions

This paper addressed several persistent gaps in the design, study, and evaluation of LLM-based pedagogical agents through Copa, a collaborative peer agent grounded in established learning theory that integrates adaptive scaffolding, multimodal learner modeling, bounded autonomy, and interpretable feedback in a single system. We paired this design contribution with a robust classroom evaluation: rather than relying on a single metric, lab study, or precompiled benchmark, we studied Copa in an authentic high-school setting across multiple lenses, including adaptivity, conceptual understanding, dependence, confidence, interpretability, and longitudinal interaction patterns. We found that Copa

adapted its scaffolding as students progressed, supported students in articulating conceptual understanding, and helped build confidence without fostering over-reliance. Its multi-agent, multimodal design also enabled interpretable feedback grounded in student input data and revealed distinct longitudinal engagement patterns through which students achieved similar learning gains. Together, this work connects technical novelty with pedagogical benefit, evidencing how theory-aligned LLM agents can be designed, evaluated, and deployed in classrooms.

Our findings offer several implications for theory, methodology, and practice, particularly for the human side of classroom AI integration. Social Constructivism and SCT frame learning as socially mediated, thereby foregrounding the importance of agents' social presence in LLM pedagogical agents (Xu, Lin, Gorter, Schneider, Weidlich, Davis, Kreijns and de Groot, 2026; Dever, Wiedbusch, Romero and Azevedo, 2024; Lyu, Li, Li, Oh, Song, Zhu and Xing, 2025). As a conversational partner, the agent invariably exerts social influence on students' reasoning, responses, and collaboration. When students ignored Copa's feedback and became frustrated, they often articulated and negotiated their understanding with each other. Thus, *agents can function as instigators for collaborative sense-making even when students disengage from them directly*. Frustration, though often viewed negatively, can be productive when it results in collaboration and critical thinking, as long as it does not cause disengagement or demoralization. This aligns with *productive failure* perspectives showing that struggle can support learning, so long as students remain engaged (Chowrira, Smith, Dubois and Roll, 2019; Sinha and Kapur, 2021; Kapur and Bielaczyc, 2012; Cohn, Snyder, Fonteles, T S, Montenegro and Biswas, 2024b).

Our prior work describes the "helpfulness paradox": agents must be helpful enough to invite engagement, yet restrained enough to prevent cognitive offloading (Cohn et al., 2026). Our results extend this idea by showing that agents can support productive learning and collaborative behavior even when students ignore their guidance. When students *did* act on Copa's feedback, these episodes shaped later interaction: seeing that Copa's prompts could yield tangible progress, students became more willing to engage with its feedback. This balance may explain both the lack of over-reliance and the limits of trust we observed: Copa was useful enough for students to seek validation or help when stuck, but not so authoritative that they deferred to it when they believed they could solve the problem themselves.

While prior work has emphasized student agency in interactions with intelligent systems (Xing, Kim, Song, Li, Li and Kim, 2026; Yan, Pammer-Schindler, Mills, Nguyen and Gašević, 2025), agency alone is insufficient. Students must also *buy in* to the agent's value to use it productively, yet this dimension remains underexplored. Such buy-in requires students to view pedagogical AI systems as *collaborators* rather than *answer providers*, and that hybrid human-AI systems are designed in ways that are understandable to students. Progress on this front will require collaboration across disciplines and stakeholders, a view echoed in recent literature (Alfredo et al., 2024; Saqr, Misiejuk and López-Pernas, 2026; Bo, 2025; Järvelä, Nguyen and Hadwin, 2023; Cohn et al., 2024b; Cohn, Davalos, Vatrál, Fonteles, Wang, Ma and Biswas, 2024a).

We provide the following limitations in the interest of transparency:

- Because institutional constraints precluded an RCT, we make correlational rather than causal claims; however, authentic, non-randomized classroom studies often provide meaningful evidence in *Computers & Education* research (e.g., via clustering; Juan, Lee and Wu, 2026).
- Evaluation targets within-system behavior rather than direct comparison with others.
- Researcher-led classroom instruction may limit generalizability to independent teacher implementation.
- Self-regulatory and longitudinal interaction patterns were analyzed qualitatively or through case studies, requiring deeper quantitative validation.
- The boundary between productive and unproductive frustration was not explicitly operationalized or evaluated.
- Student-agent interaction volume was modest, despite reflecting authentic classroom use.
- Confidence was inferred from dialogue-state proxies rather than triangulated through student self-reports.
- Interpretability claims rely on lexical- and embedding-based analyses and require further stakeholder validation.

Future work will address these limitations by conducting larger-scale classroom studies with independent teacher implementation, incorporating comparison conditions and ablation analyses to better isolate the contributions of

Copa's architectural components, and extending deployments over longer periods of time. We will also deepen the quantitative analysis of self-regulatory and longitudinal interaction patterns, explicitly investigate the boundary between productive and unproductive frustration, and triangulate confidence-related dialogue markers with student self-report measures. Finally, we will examine the practical value of Copa's interpretability through stakeholder-centered studies with teachers and students, focusing on whether the agent's reasoning artifacts meaningfully support understanding, auditing, and instructional decision-making.

Taken together, these directions aim to strengthen the empirical, methodological, and practical foundations needed to design pedagogical agents that are both educationally effective and responsibly deployed in classrooms. We hope this work inspires subsequent pedagogical AI systems and research that support learners without fostering dependence, balancing guidance with autonomy to promote critical thinking and durable learning.

Declaration of Competing Interest

The authors report no competing financial or personal interests that could be perceived as having influenced the research, its interpretation, or the reporting of its results.

Declaration of Generative AI in the Writing Process

The manuscript text was prepared by the authors and reflects their own interpretation of the study, analyses, and findings. ChatGPT was used during drafting to support improvements in language and readability. Any AI-assisted wording was subsequently reviewed and revised by the authors to ensure that the intended meaning and scholarly claims were preserved. Responsibility for the final content rests entirely with the authors.

References

- Alfredo, R., Echeverria, V., Jin, Y., Yan, L., Swiecki, Z., Gašević, D., Martinez-Maldonado, R., 2024. Human-centred learning analytics and ai in education: A systematic literature review. *Computers and Education: Artificial Intelligence* 6, 100215.
- Aulia, F.E., Hidayah, I., Fauziati, S., Anissa, S., 2025. Guiding self-regulated learning with an llm-based pedagogical chatbot, in: 2025 International Seminar on Application for Technology of Information and Communication (iSemantic), IEEE. pp. 505–511.
- Baker, R.S., Corbett, A.T., Koedinger, K.R., Wagner, A.Z., 2004. Off-task behavior in the cognitive tutor classroom: When students "game the system", in: Proceedings of the SIGCHI conference on Human factors in computing systems, pp. 383–390.
- Bandura, A., 2001. Social cognitive theory: An agentic perspective. *Annual review of psychology* 52, 1–26.
- Birks, M., Chapman, Y., Francis, K., 2008. Memoing in qualitative research: Probing data and processes. *Journal of research in nursing* 13, 68–75.
- Bo, N.S.W., 2025. Oecd digital education outlook 2023: Towards an effective education ecosystem. *Hungarian Educational Research Journal* 15, 284–289.
- Chowrira, S.G., Smith, K.M., Dubois, P.J., Roll, I., 2019. Diy productive failure: boosting performance in a large undergraduate biology course. *npj Science of Learning* 4, 1.
- Chu, Y., Li, H., Yang, K., Shomer, H., Liu, H., Copur-Gencetuk, Y., Tang, J., 2025. A LLM-Powered Automatic Grading Framework with Human-Level Guidelines Optimization, in: Educational Data Mining, International Educational Data Mining Society. p. n/a. URL: <https://educationaldatamining.org/EDM2025/proceedings/2025.EDM.long-papers.80/index.html>, doi:10.48550/arXiv.2410.02165.
- Cock, J.M., Marras, M., Giang, C., Käser, T., 2022. Generalisable methods for early prediction in interactive simulations for education, in: Mitrovic, A., Bosch, N. (Eds.), Proceedings of the 15th International Conference on Educational Data Mining, International Educational Data Mining Society, Durham, United Kingdom. pp. 183–194. doi:10.5281/zenodo.6852968.
- Cohn, C., Davalos, E., Vatrál, C., Fonteles, J.H., Wang, H.D., Ma, M., Biswas, G., 2024a. Multimodal methods for analyzing learning and training environments: A systematic literature review. Submitted to ACM Computing Surveys. Currently under review URL: <https://arxiv.org/abs/2408.14491>.
- Cohn, C., Fonteles, J.H., Snyder, C., Srivastava, N., T S, A., Campbell, D., Montenegro, J., Biswas, G., 2025a. Exploring the design of pedagogical agent roles in collaborative stem+c learning, in: Proceedings of the 18th International Conference on Computer-Supported Collaborative Learning-CSCL 2025, pp. 330–334, International Society of the Learning Sciences.
- Cohn, C., Guo, S., Rayala, S., Wang, H.D., Mohammed, N., Timalisina, U., Jain, S., Eeds, A., Deweese, M., Osborn Popp, P.J., Stanton, R., Walker, S., Ma, M., Biswas, G., 2026. Evidence-decision-feedback: Theory-driven adaptive scaffolding for llm agents. Submitted to the 27th International Conference on Artificial Intelligence in Education (AIED). Currently under review .
- Cohn, C., Rayala, S., Snyder, C., Fonteles, J.H., Jain, S., Mohammed, N., Timalisina, U., Burriss, S., T S, A., Srivastava, N., Deweese, M., Eeds, A., Biswas, G., 2025b. Personalizing student-agent interactions using log-contextualized retrieval augmented generation (rag), in: International Conference on Artificial Intelligence in Education Workshop on Epistemics and Decision-Making in AI-Supported Education., Springer. URL: <https://arxiv.org/abs/2505.17238>.
- Cohn, C., Rayala, S., Srivastava, N., Fonteles, J.H., Jain, S., Luo, X., Mereddy, D., Mohammed, N., Biswas, G., 2025c. A theory of adaptive scaffolding for llm-based pedagogical agents, in: Accepted to the AAAI Conference on Artificial Intelligence. In press.

- Cohn, C., Snyder, C., Fonteles, J.H., T S, A., Montenegro, J., Biswas, G., 2024b. A multimodal approach to support teacher, researcher and ai collaboration in stem+ c learning environments. *British Journal of Educational Technology* .
- Cohn, C., Snyder, C., Montenegro, J., Biswas, G., 2024c. Towards a human-in-the-loop llm approach to collaborative discourse analysis, in: *International Conference on Artificial Intelligence in Education*, Springer. pp. 11–19.
- Cohn, C., T S, A., Mohammed, N., Biswas, G., 2025d. CoTAL: Human-in-the-Loop Prompt Engineering for Generalizable Formative Assessment Scoring. Submitted to the *International Journal of Artificial Intelligence in Education (IJAIED)*. Currently under review .
- Dai, L., Jiang, Y.H., Chen, Y., Guo, Z., Liu, T.Y., Shao, X., 2025. Agent4EDU: Advancing AI for Education with Agentic Workflows, in: *Proceedings of the 2024 3rd International Conference on Artificial Intelligence and Education*, Association for Computing Machinery, New York, NY, USA. pp. 180–185. URL: <https://dl.acm.org/doi/10.1145/3722237.3722268>, doi:10.1145/3722237.3722268.
- Dever, D.A., Wiedbusch, M.D., Romero, S.M., Azevedo, R., 2024. Investigating pedagogical agents' scaffolding of self-regulated learning in relation to learners' subgoals. *British Journal of Educational Technology* 55, 1290–1308.
- Di Mitri, D., Srivastava, N., Fernandez Nieto, G., Echeverria, V., Cobos, R., Tudur Sadashiva, A., Spikol, D., Wong, K.Y.C., Zhou, Q., Cukurova, M., 2025. The second international workshop on multimodal artificial intelligence in education (maied'25), in: *International Conference on Artificial Intelligence in Education*, Springer. pp. 292–299.
- Fonteles, J.H., Cohn, C., Ayalon, E., Zhou, M., T.S., A., Davalos, E., Li, Z., Rayala, S., Mereddy, D., Coursey, A., Jain, S., Zhang, Y., Enyedy, N., Danish, J., Biswas, G., 2026. Analyzing embodied learning in classroom settings: A human-in-the-loop ai approach for multimodal learning analytics. *Learning and Instruction* 103, 102274. doi:<https://doi.org/10.1016/j.learninstruc.2025.102274>.
- Ganguly, A., Mehjabin, N., Malik, A., Johri, A., 2026. Conversational ai agents in education: An umbrella review of current utilization, challenges, and future directions for ethical and responsible use. *AI and Ethics* 6, 72.
- Hakami, E., Hernández Leo, D., 2020. How are learning analytics considering the societal values of fairness, accountability, transparency and human well-being?: A literature review. Martínez-Monés A, Álvarez A, Caeiro-Rodríguez M, Dimitriadis Y, editors. *LASI-SPAIN 2020: Learning Analytics Summer Institute Spain 2020: Learning Analytics. Time for Adoption?*; 2020 Jun 15-16; Valladolid, Spain. Aachen: CEUR; 2020. p. 121-41 .
- Hatch, J.A., 2002. *Doing qualitative research in education settings*. SUNY Press.
- Hou, X., Forsyth, C., Andrews-Todd, J., Rice, J., Cai, Z., Jiang, Y., Zapata-Rivera, D., Graesser, A., 2025. An LLM-Enhanced Multi-agent Architecture for Conversation-Based Assessment, in: Cristea, A.I., Walker, E., Lu, Y., Santos, O.C., Isotani, S. (Eds.), *Artificial Intelligence in Education*, Springer Nature Switzerland, Cham. pp. 119–134. doi:10.1007/978-3-031-98417-4_9.
- Hutchins, N.M., Biswas, G., Maróti, M., Lédeczi, Á., Grover, S., Wolf, R., Blair, K.P., Chin, D., Conlin, L., Basu, S., et al., 2020. C2stem: A system for synergistic learning of physics and computational thinking. *Journal of Science Education and Technology* 29, 83–100.
- Järvelä, S., Nguyen, A., Hadwin, A., 2023. Human and artificial intelligence collaboration for socially shared regulation in learning. *British Journal of Educational Technology* 54, 1057–1076.
- Jin, L., Lin, B., Hong, M., So, H.J., Zhang, K., 2025. Learning by teaching: Enhancing music learning through llm-based teachable agents, in: *International Conference on Artificial Intelligence in Education*, Springer. pp. 148–155.
- Juan, Y.C., Lee, Y.H., Wu, J.Y., 2026. Generative artificial intelligence augments social interactivity and learning outcomes: Advancing the framework of a scaffolded human–genai shared agency. *Computers & Education* , 105564.
- Kapur, M., Bielaczyc, K., 2012. Designing for productive failure. *Journal of the Learning Sciences* 21, 45–83.
- Khosravi, H., Shum, S.B., Chen, G., Conati, C., Tsai, Y.S., Kay, J., Knight, S., Martinez-Maldonado, R., Sadiq, S., Gašević, D., 2022. Explainable artificial intelligence in education. *Computers and education: artificial intelligence* 3, 100074.
- Khosrawi-Rad, B., Keller, P.F., Benner, D., Grogorick, L., Borchers, A., Janson, A., Leimeister, J.M., Robra-Bissantz, S., 2025. Promoting students' motivation in language education with gamified pedagogical conversational agents. *Computers & Education* 238, 105374.
- Koedinger, K.R., Corbett, A.T., Perfetti, C., 2012. The knowledge-learning-instruction framework: Bridging the science-practice chasm to enhance robust student learning. *Cognitive science* 36, 757–798.
- Kosmyna, N., et al., 2025. Your brain on chatgpt: Accumulation of cognitive debt when using an ai assistant for essay writing task. *arXiv:2506.08872*
- Lajoie, S.P., Li, S., 2023. Theory-driven design of aided systems for enhanced interaction and problem-solving, in: *Handbook of artificial intelligence in education*. Edward Elgar Publishing, pp. 229–249.
- Li, Q., Xie, Y., Chakravarty, S., Lee, D., 2024. EduMAS: A Novel LLM-Powered Multi-Agent Framework for Educational Support, in: *2024 IEEE International Conference on Big Data (BigData)*, pp. 8309–8316. URL: <https://ieeexplore.ieee.org/abstract/document/10826103/authors>, doi:10.1109/BigData62323.2024.10826103. ISSN: 2573-2978.
- Liu, B., Zhang, J., Lin, F., Jia, X., Peng, M., 2025a. One size doesn't fit all: A personalized conversational tutoring agent for mathematics instruction, in: *Companion Proceedings of the ACM on Web Conference 2025*, pp. 2401–2410.
- Liu, J., Chen, T., Li, S., Xia, Y., Zhu, H., Wu, R., Cao, Q., Gong, X., Wu, L., 2025b. Llm-based pedagogical agent for icu simulation instructor training: A quasi-experimental study. *Nurse Education Today* , 106901.
- Lyu, B., Li, C., Li, H., Oh, H., Song, Y., Zhu, W., Xing, W., 2025. The role of teachable agents' personality traits on student-ai interactions and math learning. *Computers & Education* 234, 105314.
- Memarian, B., Doleck, T., 2024. Human-in-the-loop in artificial intelligence in education: A review and entity-relationship (er) analysis. *Computers in Human Behavior: Artificial Humans* 2, 100053.
- Mislevy, R.J., Almond, R.G., Lukas, J.F., 2003. A brief introduction to evidence-centered design. *ETS Research Report Series* 2003, i–29.
- Moos, D.C., Azevedo, R., 2009. Learning with computer-based learning environments: A literature review of computer self-efficacy. *Review of educational research* 79, 576–600.
- Munshi, A., Biswas, G., Baker, R., Ocumpaugh, J., Hutt, S., Paquette, L., 2023. Analysing adaptive scaffolds that help students develop self-regulated learning behaviours. *Journal of Computer Assisted Learning* 39, 351–368.

- Office of Educational Technology, 2023. Artificial intelligence and the future of teaching and learning: Insights and recommendations. US Department of Education, Office of Educational Technology.
- Ouyang, F., Bai, X., 2026. A systematic review of multimodal learning analytics in computer-supported collaborative learning. *Computers & Education* , 105574.
- Pan, M.Z., Arabzadeh, N., et al., 2025. Measuring agents in production. *arXiv preprint arXiv:2512.04123* .
- Ponton, M.K., Rhea, N.E., 2006. Autonomous learning from a social cognitive perspective. *New Horizons in Adult Education and Human Resource Development* 20, 38–49.
- Ritter, F.E., Tehranchi, F., Oury, J.D., 2019. Act-r: A cognitive architecture for modeling cognition. *Wiley Interdisciplinary Reviews: Cognitive Science* 10, e1488.
- Saqr, M., Misiejuk, K., López-Pernas, S., 2026. Human-ai collaboration or obedient and often clueless ai in instruct, serve, repeat dynamics? *The Internet and Higher Education* 70, 101087. URL: <https://www.sciencedirect.com/science/article/pii/S109675162600014X>, doi:<https://doi.org/10.1016/j.iheduc.2026.101087>.
- Scholz, N., Nguyen, M.H., Singla, A., Nagashima, T., 2025. Partnering with ai: A pedagogical feedback system for llm integration into programming education, in: *European Conference on Technology Enhanced Learning*. Springer. pp. 243–248.
- Schunk, D.H., 1994. Motivating self-regulation of learning: The role of performance attributions.
- Schunk, D.H., Usher, E.L., 2012. Social cognitive theory and motivation. *The Oxford handbook of human motivation* 2, 11–26.
- Shi, Y., Liang, R., Xu, Y., 2025. EducationQ: Evaluating LLMs' Teaching Capabilities Through Multi-Agent Dialogue Framework, in: Che, W., Nabende, J., Shutova, E., Pilehvar, M.T. (Eds.), *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Association for Computational Linguistics, Vienna, Austria. pp. 32799–32828. URL: <https://aclanthology.org/2025.acl-long.1576/>, doi:10.18653/v1/2025.acl-long.1576.
- Shute, V.J., 2011. Stealth assessment in computer-based games to support learning. *Computer games and instruction* 55, 503–524.
- Sinha, T., Kapur, M., 2021. When problem solving followed by instruction works: Evidence for productive failure. *Review of Educational Research* 91, 761–798.
- Sixu, A., Yu, Y., Yungsi, M., Guandong, X., et al., 2024. Developing an llm-empowered agent to enhance student collaborative learning through group discussion, in: *International Conference on Computers in Education*, p. n/a.
- Snyder, C., 2024. Understanding Students' Collaborative Problem Solving during STEM+ C Learning using Multimodal Analysis. Ph.D. thesis. Vanderbilt University.
- Snyder, C., Hutchins, N.M., Cohn, C., Fonteles, J.H., Biswas, G., 2024. Analyzing students collaborative problem-solving behaviors in synergistic stem+ c learning, in: *Proceedings of the 14th Learning Analytics and Knowledge Conference*, pp. 540–550.
- Stamper, J., Xiao, R., Hou, X., 2024. Enhancing llm-based feedback: Insights from intelligent tutoring systems and the learning sciences, in: *International Conference on Artificial Intelligence in Education*, Springer. pp. 32–43.
- Stryker, C., 2025. What is agentic ai? URL: <https://www.ibm.com/think/topics/agentic-ai>.
- Sun, E., Tai, L., 2025. MultiTutor: Collaborative LLM Agents for Multimodal Student Support, in: *Proceedings of the Innovation and Responsibility in AI-Supported Education Workshop*, PMLR. pp. 174–190. URL: <https://proceedings.mlr.press/v273/sun25a.html>. ISSN: 2640-3498.
- Timalsina, U., Davalos_Anaya, E., Sanda, N., Zhang, Y., Horn_Fonteles, J., T_S, A., Biswas, G., 2025. Syncflow: A scalable platform for multimodal learning analytics, in: *US Research Software Engineering Conference (USRSE25)*, Zenodo. p. n/a.
- Tsai, Y.S., Whitelock-Wainwright, A., Gašević, D., 2021. More than figures on your laptop:(dis) trustful implementation of learning analytics. *Journal of Learning Analytics* 8, 81–100.
- Vatral, C., Biswas, G., Goldberg, B., 2023. A theoretical framework for multimodal learner modeling and performance analysis in experiential learning environments., in: *AI-GEL@ AIED*, pp. 68–78.
- Vygotsky, L.S., 1978. *Mind in society: The development of higher psychological processes*. volume 86. Harvard university press.
- Wang, T., Zhan, Y., Lian, J., Hu, Z., Yuan, N.J., Zhang, Q., Xie, X., Xiong, H., 2025. LLM-powered Multi-agent Framework for Goal-oriented Learning in Intelligent Tutoring System, in: *Companion Proceedings of the ACM on Web Conference 2025*, Association for Computing Machinery, New York, NY, USA. pp. 510–519. URL: <https://dl.acm.org/doi/10.1145/3701716.3715244>, doi:10.1145/3701716.3715244.
- Wu, T., Zhai, X., Song, Y., 2025. The effects of gai-enhanced pedagogical agents in the metaverse (gpaim) on elementary school students' conceptual understanding and cognitive engagement patterns. *Computers & Education* , 105555.
- Xi, L., Zhang, Y., Wang, Q., 2025. Investigating the effects of an llm-based socratic conversational agent on students' academic performance and reflective thinking in higher education. *Computers & Education* , 105494.
- Xing, W., Kim, T., Song, Y., Li, H., Li, C., Kim, J., 2026. Unveiling interaction patterns between students and generative ai teachable agent: Focusing on students' agency and ai agents' authority. *British Journal of Educational Technology* .
- Xu, K.M., Lin, L., Gorter, M., Schneider, S., Weidlich, J., Davis, R.O., Kreijns, K., de Groot, R., 2026. Social presence: A key factor in embedding a pedagogical agent into online learning in primary education. *British Journal of Educational Technology* 57, 227–242.
- Yan, L., Pammer-Schindler, V., Mills, C., Nguyen, A., Gašević, D., 2025. Beyond efficiency: Empirical insights on generative ai's impact on cognition, metacognition and epistemic agency in learning. *British Journal of Educational Technology* 56, 1675–1685.
- Zha, S., Liu, Y., Zheng, C., Xu, J., Yu, F., Gong, J., Xu, Y., 2025. Mentigo: An Intelligent Agent for Mentoring Students in the Creative Problem Solving Process, in: *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, Association for Computing Machinery, New York, NY, USA. pp. 1–22. URL: <https://dl.acm.org/doi/10.1145/3706598.3713952>, doi:10.1145/3706598.3713952.
- Zhang, J., Borchers, C., Cohn, C., Srivastava, N., Snyder, C., Guo, S., TS, A., Mohammed, N., Noh, H., Biswas, G., 2026. Using large language models to detect socially shared regulation of collaborative learning, in: *Proceedings of the LAK26: 16th International Learning Analytics and Knowledge Conference*, pp. 883–890.

- Zhang, N., Biswas, G., McElhaney, K.W., Basu, S., McBride, E., Chiu, J.L., 2020. Studying the interactions between science, engineering, and computational thinking in a learning-by-modeling environment, in: International conference on artificial intelligence in education, Springer. pp. 598–609.
- Zhang, X., Zhang, C., Sun, J., Xiao, J., Yang, Y., Luo, Y., 2025. Eduplanner: Llm-based multi-agent systems for customized and intelligent instructional design. IEEE Transactions on Learning Technologies .
- Zhou, Y., Pankiewicz, M., et al., 2025. Impact of llm feedback on learner persistence in programming, in: International Conference on Computers in Education, p. n/a.