

MATH-AI: An Integrated Conceptual Framework for Student-AI (LLM) Interactions

Katya Hernández Holliday
San Diego State University

With the rising popularity of generative artificial intelligence, particularly large language models (LLMs), it is imperative to understand how students are leveraging these rapidly evolving tools in their learning. Although researchers have begun to address issues of academic integrity, acceptance, and adoption, few conceptual frameworks exist for analyzing how students interact with LLMs in the context of mathematics problem solving. This paper introduces the MATH-AI framework, a synthesis of theories of metacognition, student agency, trust in automation, and human processing and regulation to examine how AI mediates math problem solving and potentially reconfigures the learning process. A case study of two students is provided to demonstrate the framework's application and illustrate distinct patterns of trust, processing, and self-regulation during AI-mediated mathematics problem solving. These cases highlight the potential risks and opportunities associated with LLM use and provide a foundation for future empirical studies and development of subject-specific AI guidelines.

Keywords: Undergraduate Mathematics, LLM, Artificial Intelligence, ChatGPT, Learning

Despite rapid adoption and integration of generative artificial intelligence (GenAI), especially large language models (LLMs), within higher education, little is known about how these tools shape cognitive processes during math problem solving. Influenced by institutional promotion and social normalization (Cialdini et al., 1991), many have come to view adoption of LLMs (e.g., ChatGPT, Google Gemini) as inevitable. Nonetheless, some worry about how prolonged use may interfere with skill development and learning (Nguyen et al., 2025).

While studies have addressed human-AI interactions within the university context, they often focus on issues of acceptance and adoption (Ivanov et al., 2024; Xu et al., 2025), academic integrity (Almisad & Aleidan, 2025), and perceptions among students and faculty (Tsiani et al., 2025). These studies provide useful insights about how individuals make decisions about how to use GenAI but do not address how such tools reconfigure student thinking and reasoning.

This paper proposes a framework for analyzing student-AI interactions with the aim of revealing conditions that support cognitive engagement and processes essential for learning math. While models of AI-human interaction typically center factors that influence adoption, perceived utility, or ethics, this framework highlights foundational academic practices such as studying, reasoning, and ultimately learning. We draw on frameworks from cognitive science (i.e., cognitive processing), educational psychology (i.e., self-efficacy, self-regulated learning), and human-computer dynamics (i.e., trust) to develop an interdisciplinary model for examining student-AI interactions during mathematics problem solving. By applying an interdisciplinary, empirically guided approach, this model serves as a conceptual tool for examining student-AI interactions and as a foundation for subject-specific, evidence-based AI use guidelines.

Theoretical Framework

Student Beliefs & Technology Adoption

Students who believe in their ability to learn often work harder, persist longer, and perform better academically (Bandura, 1997). At the same time, students hold different beliefs about their

abilities according to domain (e.g., math abilities vs. writing abilities). On a smaller scale, they gauge task-specific *self-efficacy*: whether they believe they can perform specific actions related to a task (Bandura, 1986; Pajares & Miller, 1995). Students with low *math* self-efficacy (Li et al., 2021) often exhibit higher levels of math anxiety, negatively impacting performance (Beilock & Maloney, 2015). In contrast, feelings of competence and autonomy have been linked to cognitive activation (e.g., attention, engagement, and thinking; Lu et al., 2022). Researchers have noted the important role teachers play in supporting students' task- and domain-specific self-efficacy by providing targeted feedback and guidance (Lu et al., 2022; Van der Beek et al., 2017). Open questions remain whether students prompt LLMs like ChatGPT in ways that result in such support, and whether they act on generated feedback during math problem solving.

Apart from beliefs about their own abilities, students' beliefs about chatbots' abilities may also influence use. Research on attitudes towards technology adoption suggests that positive attitudes often predict adoption and use (Ajzen, 1991). In the case of ChatGPT, students who hold positive attitudes extract more benefits by asking more detailed questions or experimenting with prompting strategies, leading to more productive uses than those who remain skeptical (Zhuang, 2025). However, students' trust in generated output can also enable overreliance. Theories of trust in automation describe over trust as a form of miscalibration whereby users rely on an agent to perform tasks beyond its capabilities. On the other hand, distrust often leads to underuse or disuse (Lee & See, 2004). Trust, therefore, guides reliance; a student who over trusts an LLM may override their own thinking/reasoning, uncritically adopting inaccurate output while students who avoid the tool may miss out on opportunities for support.

Cognitive Processing, Offloading, & Self-Regulation

Students' perceptions and beliefs manifest in how they engage both with concepts and tools during problem solving. Prompting LLMs may reduce memory demands by providing immediate information. Decreasing the working memory requirements of a task can reduce cognitive load (Paas et al., 2003), but the benefits of cognitive offloading (reducing cognitive demand by externalizing some part of a task) are contingent on whether offloaded demands are required to realize the learning goal (Risko & Gilbert, 2016). Offloading *mastered* procedures (e.g., arithmetic, symbol manipulation) may free cognitive capacity for conceptual reasoning, whereas offloading reasoning processes central to the learning goal may threaten learning.

To investigate when offloading supports or undermines learning, it is necessary to consider the types of cognitive processing students engage in during AI interactions. Craik and Lockhart (1972) differentiate shallow from deep processing where shallow processing consists of surface-level encoding (e.g., tending to appearance or sound), typically resulting only in short-term retention. In contrast, deep processing involves creating associations, facilitating long-term retention (McLeod, 2025). Between shallow and deep processing, the ongoing and sometimes incomplete process of self-explaining promotes integration of new information (Chi et al., 1994).

In addition to cognitive processing, students' capacity for self-regulation further shapes how they engage with LLMs. Self-regulated learning involves setting goals, monitoring one's own understanding, and adapting strategies in response to feedback (Zimmerman, 2000). However, unlike human instructors, LLMs do not provide formative feedback on students' progress toward conceptual understanding unless they are prompted to do so. Students must also monitor their progress while tending to the accuracy of output. Strong self-regulation skills can reduce uncritical acceptance of LLM output, while the lack of such skills may leave students vulnerable to over trusting tools and/or disengaging from problem solving altogether.

Practicing problem solving is a key aspect of studying mathematics. Studying, a key form of self-regulated learning, usually occurs without direct expert/teacher guidance, requiring students to rely on their own strategies or collaborate with peers. It often involves gathering and synthesizing information and usually results in the production of observable indicators of cognitive processing such as notes, annotations, or summaries (Winne & Hadwin, 1998). LLMs introduce a new dynamic; these tools can generate similar outputs, potentially reconfiguring how students approach studying. To understand this shift, it is important to examine the self-regulation strategies students employ when working with LLMs.

Potential Outcomes of AI Reliance During Math Problem Solving

Beliefs, patterns of use, and self-regulation determine the outcomes of students' interactions with LLMs. Depending on the quality of cognitive engagement, AI can support learning or pose significant risks including metacognitive miscalibration, changes in overall self-efficacy, and entrenchment of robust misconceptions. Recent reports (OpenAI, 2025) announce that ChatGPT has shown its strongest performance on academic tasks to date. As LLMs improve and grow in popularity, students may begin to treat them much like the internet, as general-purpose stores of knowledge. Studies have shown that internet searching can create the illusion of knowledge as individuals mistake access to information for a reflection of personal understanding, even in the absence of understanding (Fisher et al., 2015). These experiences may be further exacerbated by *metacognitive miscalibrations* where students overestimate their own knowledge (Lichtenstein & Fischhoff, 1977), and by new AI-specific challenges where students overestimate tool accuracy.

AI systems are designed to promote engagement, not challenge users (unless prompted to do so) and are more likely to reinforce assumptions embedded within prompts. Thus, an LLM is more likely to enable robust misconceptions if students interpret output as confirmation of understanding. Robust misconceptions are especially resistant to change and may become further entrenched rather than corrected in the absence of opposition (Smith et al., 1993).

Persistent reliance on LLMs may also influence math self-efficacy. If miscalibrations lead to overconfidence, lack of knowledge may be made apparent through formal assessments and academic performance. Furthermore, shallow forms of processing may become habitual through repetition. These risks can negatively affect students' learning outcomes and even their agency as they come to depend on LLM tools in ways that limit their abilities to regulate and sustain their own learning (Krook, 2025). Taken together, these constructs highlight the complex ways in which LLMs may interfere with cognitive processes and subsequent learning outcomes, providing a foundation to theorize student-AI interactions as part of the math learning process.

Conceptual Framework

As GenAI becomes increasingly present in educational contexts, we must clarify the cognitive and social processes students engage with when working with AI and identify frameworks that help math educators integrate these tools without undermining learning. Guided by these goals, the present framework emerged as a response to observations of students interacting with AI in ways that reflected less-than-optimal use. This framework is empirically informed by recurring patterns observed in interview data with undergraduates solving algebra problems with ChatGPT. Students in this study often operated in a space between having no strategy and having some relevant prior knowledge that *could* be activated. In these moments, ChatGPT enabled them to pursue solution pathways and eventually produce answers. However, interactions rarely resulted in advancements in conceptual understanding. Many students accepted output without fully integrating it with their own mathematical thinking. These

empirical insights illuminated the need for constructs that highlight levels of cognitive processing during AI interaction. They also underscored the need for interdisciplinary knowledge that captured students' self-perceptions, metacognitive monitoring, and sociocultural perspectives on tool use to describe the complex ways in which students navigate AI assistance/availability and their own cognitive faculties.

To capture the complex interplay between students' beliefs and perceptions, cognitive processes, and how these manifest in LLM use patterns, we propose the MATH-AI framework which encompasses five key dimensions: *metacognition*, *agency*, *trust*, *human processing/regulation*, and *AI mediation*. This framework integrates established constructs to address emerging challenges in mathematics education and offers a structure for educators and researchers to analyze use patterns and link them to potential benefits and risks. The five dimensions of MATH-AI are conceptualized as follows:

1. **Metacognition** refers to students' abilities to monitor and evaluate their own knowledge and performance. It captures how students assess their understanding during AI interaction, whether monitoring supports or undermines cognitive engagement.
2. **Agency** refers to students' sense of ownership/control over their learning. It is shaped by beliefs about their ability to learn and do math (self-efficacy) as well as their willingness and motivation to engage in the learning process. It captures students' orientations toward using AI (as a scaffold or a crutch). This is a dispositional construct that highlights whether students position themselves as active rather than passive in their engagement. Research has established connections between self-efficacy and motivation, yet little is known about how these are influenced by AI interaction. Early findings suggest a tendency of students to depend on AI, diminishing their agency (Darvishi et al., 2024).
3. **Trust** encompasses students' judgements about output reliability. Particularly, whether trust is appropriately calibrated to specific tasks or if it enables overreliance.
4. **Human Processing & Regulation** captures the cognitive/metacognitive moves students make during problem solving including the *depth* of cognitive engagement and the *regulatory* actions taken to monitor, evaluate, and adapt strategies (e.g., interrogating/verifying/seeking explanations). While agency reflects willingness and stance toward learning, processing and regulation describe enacted behaviors and strategies resulting from that stance. This construct draws attention to how students position AI in practice, as either a collaborator/scaffold or a shortcut to bypass deeper processing.
5. **AI Mediation** operates through all prior dimensions of the framework, highlighting the ways AI tools intervene in learning processes. For example, by offering scaffolds, feedback, or alternative representations or by encouraging reliance that diminishes agency and depth of processing.

These five dimensions are operationalized across three phases: input, interaction, and output. Inputs include students' preconceptions, beliefs and internalized views. Interactions encompass how students engage with ChatGPT, including prompting and verbalized processing. Outputs capture potential consequences of sustained patterns of use, including habitual cognitive processes, impacts on agency, and the risks/opportunities for learning. Each of the phases also incorporates a feedback mechanism consisting of students' reflections during and after AI use.

Students' self-efficacy (*input*) largely contributes to persistence, efforts to seek understanding, and the types of follow up questions they formulate. When working with AI, students with higher domain- and task-specific efficacy may be less likely to cede important

problem-solving actions to an LLM. Similarly, perceptions of tool efficacy influence attitudes, trust, and choices about whether and how to use such tools during math problem solving.

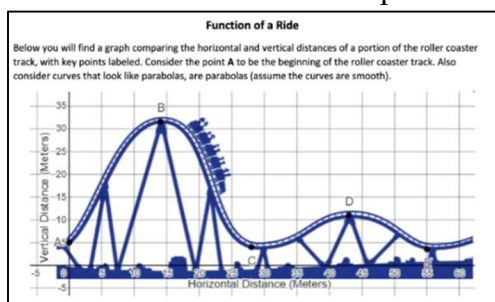
During *AI interaction*, students make decisions (perhaps subconsciously) about the depth of cognitive processing they engage in and their level of reliance. Importantly, a student may exhibit high levels of *productive* reliance. For example, using an LLM to generate practice problems, critique their reasoning, or explanations can promote self-regulation by helping students monitor their knowledge and discuss misconceptions. In contrast, reliance that offloads reasoning without verification may stall learning.

Sustained patterns of use may develop into habits that either support or threaten learning (*output* phase). Outcomes depend on prompting strategies and maintaining effort during problem solving. Over time, these outcomes can have lasting effects on students' broader sense of self-efficacy, metacognitive abilities, and ultimately student agency.

MATH-AI Framework Application

To illustrate application, we draw from a subset of data collected as part of a larger project investigating if/how students engage in productive struggle when working with ChatGPT. This subset includes undergraduates who completed an algebra task, Figure 1 during think-aloud interviews. Algebra was selected because it is a subject most undergraduates share prior knowledge/course experience in, unlike the other tasks involving statistics and calculus. During interviews, researchers encouraged students to share their thinking.

We present two contrasting cases not meant to represent generalizable findings but meant as illustrative applications of the framework and as a starting point for the development of subject-specific guidelines. Analyses are organized according to process order (input → interactions → outputs) to show how student beliefs may influence the ways they interact with ChatGPT and how these interactions can lead to specific outcomes.



(a)

1. What is the domain and range of the function?
2. Find the intervals where the function is decreasing and increasing. How did you know?
3. Where would you scream? Describe your ride as you travel the roller coaster. Include in your description your trip from point to point, whether you are moving up or down, and discuss what is happening to your speed.

(b)

Figure 1. Image Provided with the Roller Coaster Task (a) and the first 3 questions included in the task (b).

Case 1: High Trust, Shallow Processing

Despite expressing reservations, Participant 1 (P1), a food and nutrition major, described high behavioral reliance across coursework and a positive attitude toward ChatGPT for math problem solving: “I started using it because I saw how easy it was. You just put in the question ...my friends were like, ‘you don’t use ChatGPT? Why?’ I was just like, I don’t know, I just don’t like the cheating of it.” P1’s experiences are of particular interest because they show how LLM tools may be perceived as detrimental while simultaneously being promoted through social and institutional norms: “It made me lazy, I feel dumb. I don’t feel like I’m learning...I don’t think it’s fair that [the institution] is pushing it on students...how do you expect students not to cheat if you’re giving us ChatGPT?” P1’s problem solving involved reading tasks, assessing their own knowledge, and prompting ChatGPT whenever they perceived a gap.

Inputs. P1 exhibited low self-efficacy both in general, as evidenced above, and across tasks: “I don’t even know how to solve this...I haven’t heard domain and range since high school so I’m asking ChatGPT to explain.” Their positive attitude towards ChatGPT for math problem solving enabled a high level of trust, evidenced by unquestioned acceptance of generated outputs. When asked to explain their solution, they replied: “[ChatGPT] is giving me an example, the domain is -2 to 4 and it’s saying that it goes upwards...so in this case it would be $x = 5$ because that’s where it starts and then it ends at $x = 4$, in their case that’s where it would end, and in my case I would say... $x = 10$ so I believe my domain would be that.”

Interactions. Subsequently, P1 processed output only at a shallow level: “I’m just using their example to connect it to mine as much as I can.” P1 did not monitor their own understanding or knowledge acquisition and ultimately did not identify and fix errors: “I honestly think it did a good job at [explaining domain and range] ...I still don’t fully understand...but it did help me understand just to solve the question.” This reflection suggests that P1 did not actively participate in solving and instead ceded the process to ChatGPT. Poor performance coupled with high confidence in generated output resulted in miscalibration as P1 expressed high confidence in the solutions they provided, shown in Figure 2a, despite their inaccuracy.

Potential Outcomes. Sustained use of LLMs which follow this interaction pattern can lead to substantial detriment as users may repeatedly practice incorrect procedures that further entrench misconceptions and unproductive habits. In this case, shallow processing resulted in incorrect solutions and misconceptions without P1’s awareness. More concerning than the incorrect answers, however, was the lack of evidence of developing conceptual understanding on topics the participant clearly expressed confusion about before interaction. P1 continued to struggle with domain and range after interacting with ChatGPT, despite having encountered these concepts in high school. This pattern illustrates how high trust combined with shallow processing can lead to miscalibration, potentially impeding the development of mathematical competence while also eroding students’ broader sense of agency in their learning.

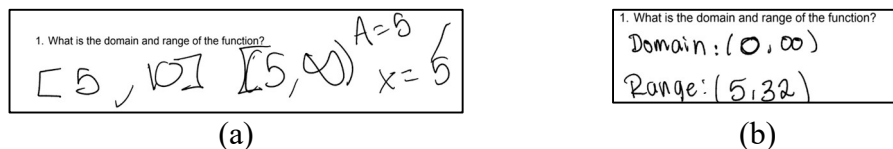


Figure 2. Written Answer to Question 1 Provided by Participant 1 (a) and Participant 2 (b).

Case 2: High Trust, Deep Processing

Participant 2 (P2), a business management major, reflected on how ChatGPT has changed their approach to mathematics: “I remember, if I needed to do a math problem, looking it up on the internet and finding places where it didn’t cost money...it would just be hard. I would always have to figure it out on my own.” They also described how sustained AI reliance seemed to undermine their learning: “I’m not doing very well in the [math] class...at the beginning of the year, when I didn’t understand stuff...I would just put it into ChatGPT...but now I’m trying a lot harder to do it on my own, so that I understand it for the final.”

Inputs. P2 described low self-efficacy and high reliance. They also expressed efforts to shift away from dependence and toward building understanding. P2 described a high level of trust in ChatGPT, noting how easy it was to complete assignments by prompting the tool.

Interactions. P2 prompted ChatGPT for conceptual help, reflecting: “I always have to remember domain is this way [traces a horizontal plane] and range is this way [traces a vertical plane],” before prompting: “how to find domain and range.” However, ChatGPT’s lengthy,

decontextualized output caused frustration: “That makes no sense...I don’t know how that helps me find the domain.” Rather than abandoning the tool, P2 re-prompted with, “when looking at a graph with porabalas [sic].” P2 reflected on output: “this made me think about how to find it with numbers, so I said by actually looking at a graph” and exhibited deep processing, making connections both to real-world contexts and prior math knowledge:

This [the right of the graph] has to connect with this [left end of the graph] eventually, right? It’s not just going to be cut off...this has to come back to point A, which makes me feel like the domain is all real numbers...but a rollercoaster isn’t infinite...obviously this is the end, so I feel like it has to stop here.

Potential Outcomes. P2 transitioned from frequent prompting to relying on their own thinking and integrated ChatGPT’s examples with their own ideas: “If it’s like the example, then the range is negative infinity to 32...but then there’s the bottom of it [points to the bottom of the graph] ...so I think it has to stop there.” Although P2 struggled to interpret output, additional prompting enabled progress toward making sense of domain and range in the situated context of a rollercoaster. P2’s written answer is shown in Figure 2b. Unlike P1, who relied on ChatGPT and shallow processing, P2 engaged critically with output. This pattern of use can support learning as students question their understanding and make connections to both prior knowledge and real-world contexts.

Discussion

The MATH-AI framework focuses on how meta/cognition, agency, trust, human processing, and AI interaction intersect to reconfigure the ways students approach common mathematics processes such as problem solving and studying. While this study was limited to two cases, it provides a foundation for empirical analyses to further extend and refine the framework, explore its applicability/comparisons across diverse student populations, and inform the design of interventions to guide students towards productive uses of LLMs for learning math. For example, a self-regulation framework informed by these analyses might prompt students to ask: Can I explain this concept without assistance? Can I solve a similar problem on my own?

To best support mathematics self-efficacy, self-regulation and cognitive processing, educators must intentionally promote productive uses of LLMs. Regardless of instructor views, AI is here to stay. The critical question now is how its mediation will shape students’ dispositions and strategies for learning mathematics. As social-cognitive theories remind us, self-efficacy develops through mastery and regulation, while models of normative behavior suggest that students’ beliefs about these tools may ultimately be overshadowed by widespread availability and adoption. At the same time, theories on cognition highlight the potential for LLM scaffolding and the risk of offloading essential cognitive processes. Together, these perspectives emphasize that the integration of AI in math education can lead to starkly different outcomes: either supporting learning or entrenching patterns of dependence that reduce students’ opportunities for deep thinking and cognitive development. By applying the MATH-AI framework to research and teaching practice, educators can gain a deeper understanding of how to support mathematical thinking and learning in an AI-powered world.

References

Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2). [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)

- Almisad, B., & Aleidan, A. (2025). Faculty perspectives on generative artificial intelligence: insights into awareness, benefits, concerns, and uses. *Frontiers in Education*, 10. <https://doi.org/10.3389/feduc.2025.1632742>
- Bandura, A. (1997). *Self-efficacy: The exercise of control*. W H Freeman/Times Books/Henry Holt & Co.
- Bandura, A., & National Institution of Mental Health. (1986). *Social foundations of thought and action: A social cognitive theory*. Prentice-Hall, Inc.
- Beilock, S. L., & Maloney, E. A. (2015) Math anxiety: A factor in math achievement not to be ignored. *Policy Insights from the Behavioral and Brain Sciences*, 2(1), 4-2. <https://doi.org/10.1177/2372732215601438>
- Chi, M. T. H., DeLeeuw, N., Chiu, M., & LaVancher, C. (1994). Eliciting self-explanations improves understanding. *Cognitive Science*, 18.
- Cialdini, R. B., Kallgren, C. A., & Reno, R. R. (1991). A focus theory of normative conduct: A theoretical refinement and reevaluation of the role of norms in human behavior. *Advances in Experimental Psychology*, 24, 201–234. [https://doi.org/10.1016/S0065-2601\(08\)60330-5](https://doi.org/10.1016/S0065-2601(08)60330-5)
- Darvishi, A., Khosravi, H., Sadiq, S., Gasevic, D., & Siemens, G. (2024). Impact of AI assistance on student agency. *Computers & Education*, 210. <https://doi.org/10.1016/j.compedu.2023.104967>
- Fisher, M., Goddu, M., & Keil, F. (2015). Searching for explanations: How the Internet inflates estimates of internal knowledge. *Journal of experimental psychology: General*, 144 (3), 674–87. <https://doi.org/10.1037/xge0000070>.
- Ivanov, S., Soliman, M., Tuomi, A., Alkathiri, N. A., & Al-Alawi, A. N. (2024). Drivers of generative AI adoption in higher education through the lens of the theory of planned behaviour. *Technology in Society*, 77. <https://doi.org/10.1016/j.techsoc.2024.102521>
- Krook, J. (2025). When autonomy breaks: the hidden existential risk of AI. *AI & Society*. <https://doi.org/10.1007/s00146-025-02397-5>
- Lee, J. D., & See, K. A. (2004). Trust in automation: designing for appropriate reliance. *Human Factors*, 46(1), 50-80. https://doi.org/10.1518/hfes.46.1.50_30392
- Li, H., Liu, J., Zhang, D., & Liu, H. (2021). Examining the relationships between cognitive activation, self-efficacy, socioeconomic status, and achievement in mathematics: A multi-level analysis. *The British Journal of Educational Psychology*, 91(1), 101-126. <https://doi.org/10.1111/bjep.12351>
- Lichtenstein, S., & Fischhoff, B. (1977). Do those who know more also know more about how much they know? *Organizational Behavior and Human Performance*, 20, 159–183.
- Lu, H., Chen, X., & Qi, C. (2022). Which is more predictive: Domain- or task-specific self-efficacy in teaching outcomes? *The British Journal of Educational Psychology*. <https://doi.org/10.1111/bjep.12554>
- McLeod, S. (2025). *Levels of processing Theory (Craik & Lockhart, 1972)*. Simply Psychology. <https://www.simplypsychology.org/levelsofprocessing.html>
- Nguyen, A., Duong, A. T., Nguyen, D. T. B., Lai, V. T. T., & Dang, B. (2025). Guidelines for learning design and assessment for generative artificial intelligence-integrated education: a unified view. *Information and Learning Sciences*. <https://doi.org/10.1108/ILS-11-2024 -0148>
- Paas, F., Renkl, A., & Sweller, J. (2003). Cognitive load theory and instructional design: Recent developments. *Educational Psychologist*, 38(1), 1–4. https://doi.org/10.1207/S15326985EP3801_1
- OpenAI. (2025). Introducing GPT-5. <https://openai.com/index/introducing-gpt-5/>

- Pajares, F., & Miller, D. M. (1995). Mathematics self-efficacy and mathematics performances: The need for specificity of assessment. *Journal of Counseling Psychology*, 42(2), 190–198. <https://doi.org/10.1037/0022-0167.42.2.190>
- Risko, E. F., & Gilbert, S. J. (2016). Cognitive Offloading. *Trends in Cognitive Sciences*, 20(9).
- Smith, J. P., diSessa, A. A., & Roschelle, J. (1993). Misconceptions reconceived: A constructivist analysis of knowledge in transition. *The Journal of the Learning Sciences*, 3(2).
- Tsiani, M., Lefkos, I., & Fachantidis, N. (2025). Perceptions of generative AI in education: Insights from undergraduate and master's-level future teachers. *Journal of Pedagogical Research*, 9(2). <https://doi.org/10.33902/JPR.202531943>
- Van der Beek, J. P. J., Van der Ven, S. H. G., Kroesbergen, E. H., & Leseman, P. P. M. (2017). Self-concept mediates the relation between achievement and emotions in mathematics. *The British Journal of Educational Psychology*. <https://doi.org/10.1111/bjep.12160>
- Winne, P. H., & Hadwin, A. F. (1998). Studying as self-regulated learning. In D. Hacker, J. Dunlosky, & A. Graesser (Eds.), *Metacognition in Educational Theory and Practice*. Erlbaum.
- Xu, J., Li, Y., Shadiev, R., & Li, C. (2025). College students' use behavior of generative AI and its influencing factors under the unified theory of acceptance and use of technology model. *Education and Information Technologies*. <https://doi.org/10.1007/s10639-025-13508-6>
- Zhuang, Y. (2025). Lessons from Using ChatGPT in Calculus: Insights from Two Contrasting Cases. *Journal of Formative Design in Learning*. <https://doi.org/10.1007/s41686-025-00098-2>
- Zimmerman, B. J. (2005). Attaining Self-regulation: A social cognitive perspective. In M. Boekaerts, P. R. Pintrich & M. Zeidner (Eds.). *Handbook of Self-Regulation*. Academic Press.