

Evaluating the Impact of LLM-guided Reflection on Learning Outcomes with Interactive AI-Generated Educational Podcasts

Vishnu Menon¹, Andy Cherney¹, Elizabeth B. Cloude², Li Zhang¹, Tiffany D. Do¹

¹Drexel University, ²Michigan State University

Correspondence: {vishnu.v.menon|harry.zhang|tiffany.do}@drexel.edu, cloudeel@msu.edu

Abstract

This study examined whether embedding LLM-guided reflection prompts in an interactive AI-generated podcast improved learning and user experience compared to a version without prompts. Thirty-six undergraduates participated, and while learning outcomes were similar across conditions, reflection prompts reduced perceived attractiveness, highlighting a call for more research on reflective interactivity design.

1 Introduction

What if educational content could not only speak to learners, but listen, adapt, interact, and assess learning processes – like *reflection* – in real-time? As learners increasingly disengage from traditional materials like textbooks (Baron and Mangen, 2021), large language models (LLMs) offer new opportunities to deliver content in more engaging, interactive, and personalized formats, such as AI-generated podcasts (Jin et al., 2025). Emerging tools like NotebookLM¹ illustrate growing public interest in generative AI for learning.

Personalized learning with AI has been shown to support self-regulated learning by encouraging learners to plan, monitor, and evaluate their progress (Shemshack and Spector, 2020; Molenaar et al., 2023). Prior work demonstrates that personalized AI-generated podcasts based on college textbooks (tailored to learners' majors, interests, and instructional preferences) can enhance learning and enjoyment compared to both textbooks and non-personalized content (Do et al., 2025). However, most AI-generated podcasts remain *passive*: learners can ask questions, but the system does not initiate interaction or assess learning to guide deeper engagement. This represents a missed opportunity, as structured interactivity has been shown to

enhance engagement and active learning in other domains (Laban et al., 2022).

More importantly, reflection is a critical component of learning – it helps learners draw meaningful and construct understanding in connection with learning goals. Embedding structured, reflection prompts into AI-generated podcasts could enhance engagement and learning, but may also disrupt learners' concentration and flow, possibly reducing their effectiveness. Design trade-offs remain unclear: when should reflection be prompted, and how can learners' responses be assessed in real-time with AI-generated podcasts?

We investigated these questions in a controlled experiment with a sample of 36 undergraduates, comparing two conditions: Reflection, where an AI-generated podcast periodically prompted learners to reflect and responded based on their input, and Standard, where no reflection prompts were prompted by the system. This study specifically investigates an AI-generated podcast featuring a single host, using two research questions: (1) Do interactive—in this case, meaning a model that can be freely interrupted, conversed with and asked questions—reflection prompts improve learning outcomes when incorporated into AI-generated podcasts compared to standard AI-generated podcasts? and, (2) Do interactive reflection prompts improve user experience when incorporated into AI-generated podcasts compared to standard AI-generated podcasts?

2 Related Work

Reflection is a key self-regulatory process that supports deeper learning and metacognitive awareness by promoting learners to contemplate their understanding and connect it with previous learning experiences. McAlpine et al. (1999) conceptualize reflection as a goal-driven process in which learners continuously integrate knowledge and action.

¹<https://notebooklm.google/>

Building on this, recent work has explored whether digital learning environments can scaffold reflection to improve learning outcomes. (Cloude et al., 2021) examined the impact of reflective prompts in a game-based learning environment with 120 adolescents. Learners received one of three types of reflection prompts during learning: progress planning, solution strategy, and different problem approaches. Findings showed that the quantity and quality of reflections influenced learning, but their effects varied depending on the learner’s goals.

Carpenter et al. (2021) further investigated reflection quality in middle-school students, using a rubric to assess written responses on a 5-point scale ranging from non-reflective to highly reflective. Higher-quality reflections – those including hypotheses, planning, and reasoning – were more predictive of learning gains. However, these reflections were scored post hoc, underscoring a key limitation in current research: the inability to evaluate and respond to reflection *in real time*. Cloude et al. (2021) highlight the lack of theoretical clarity around when and how to prompt reflection during learning, a gap this work aims to address by embedding structured, interactive prompts into AI-generated podcasts. In our system, an LLM-driven agent prompts learners to reflect and evaluates their spoken responses to guide real-time support.

While prior studies have explored personalized AI-generated podcasts for education, they have not addressed reflection. Do et al. (2025) compared generalized and personalized podcasts – generated from textbook chapters using LLMs – to traditional textbook reading on learning outcomes. Personalized podcasts, tailored to learners’ majors and interests, improved enjoyment and learning outcomes. Their systems used a multi-stage generation pipeline with Gemini 1.5 Pro to convert textbook content into conversational podcast scripts, which were then synthesized using text-to-speech models.

Other work has explored AI-generated podcasts outside of education. Yahagi et al. (2025) showed that transforming academic papers into AI-generated podcasts lowered barriers to engaging with academic literature. Similarly, Laban et al. (2022) examined AI-generated podcasts for news delivery and found that it enhanced enjoyment. However, unlike our system, these podcasts lacked interactive, reflective components and were not designed for structured learning contexts. Together, these lines of research inform our approach: we integrate structured reflection prompts – adapted in

real-time by an LLM – into AI-generated podcasts to support engagement, reflection, and learning during listening.

3 Interactive Podcast Architecture

The technical implementation of our AI-generated podcast system involves ingesting textbook content and delivering an interactive podcast on demand. The system supports two modes of interaction (Figure 1):

- *Standard*: The system delivers audio content continuously from the textbook and allows learners to interrupt at any time with questions or comments. This interaction style is similar to existing consumer podcast systems, such as NotebookLM’s Interactive mode.
- *Reflection*: The system delivers audio content and incorporates structured reflection prompts, periodically pausing after key concepts are introduced, and requiring the learner to demonstrate understanding of the content before continuing in their spoken reflection.

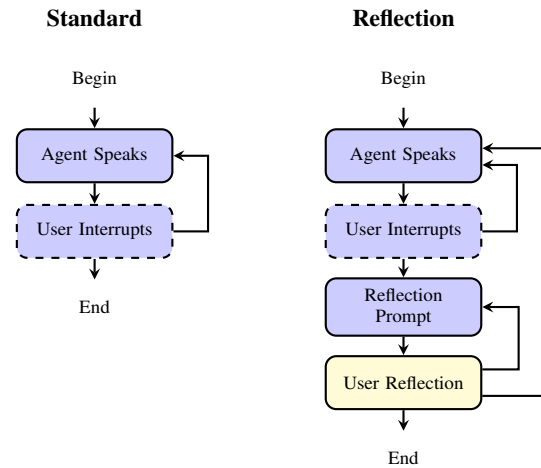


Figure 1: Standard vs Reflection Interaction Modes.

The system architecture is shown in Figure 2. The system consists of a Python backend for content generation, which hosts a LiveKit² room and creates an agent for speech synthesis. LiveKit is a platform for building AI-voice applications that can interact with users over the web.

The frontend is built with React, Next.js, and TailwindCSS, and connects to the backend to enable real-time communication with the system.

²<https://livekit.io/>

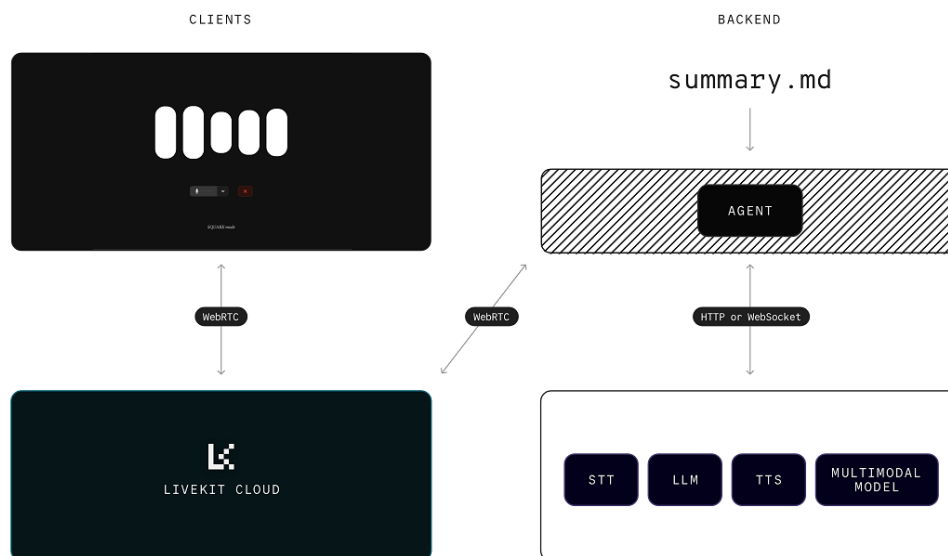


Figure 2: A diagram of our system, adapted from LiveKit’s Agents Overview.

Based on these components, our system achieves an average response latency of 300ms (d’Sa, 2024), closely matching the pace of natural human conversation (Stivers et al., 2009). This low-latency architecture, which we further detail below, enables controlled comparisons between the two interaction paradigms. To support reproducibility, we have open-sourced the system; the code repository is available on GitHub³.

3.1 Structured Summary

The system first ingests Chapter 1, Section 1.1 of OpenStax’s textbook *Introduction to Philosophy* (Smith, 2022). Then, using GPT-4 Turbo, it converts academic text into a structural and summarized skeleton for the podcast, following the summarization procedure outlined by Laban et al. (2022). We found that using Laban et al.’s structured summarization approach to generate podcasts addressed several challenges, such as material omission. When generating podcasts directly from source text, we found that the model often omitted important material and produced outputs constrained by its context window, regardless of input length. This created an artificial ceiling on content length and limited scalability for longer educational materials. By using a structured summary, where each section corresponds to a paragraph from the original source, we were able to

generate each segment independently, ensuring content coverage and improving quality, similar to skeleton-of-thought (Ning et al., 2023). The structured summary also improves interpretability, providing transparency into the generation process and facilitating easier debugging. It serves as a reference to track content coverage during the learner-facing conversation.

3.2 Podcast Generation

The structured summary is first divided into segments according to its outline and then processed by GPT-4o-mini. Each segment is used to generate corresponding portions of the podcast. Using OpenAI’s GPT-4o text-to-speech (TTS) model with the *Alloy* voice, the podcast is synthesized as natural-sounding speech, incorporating appropriate pacing and intonation based on the skeleton structure.

3.3 User Interaction and Reflection

We serve the content to the learner differently depending on their current interaction context, tracked via state machine. In the Reflection mode, it monitors learner responses to assess their knowledge of the topic and prompt reflections by asking: “So, what is the most important thing you’ve learned so far?” at the end of each section, following similar prompts by (Cloude et al., 2021). After the learner responds to the reflection prompt, we use a one-shot evaluation to determine whether the response is suitable using a binary assessment (1 =

³<https://github.com/DU-DIVALab/tutorflow>

demonstrates understanding, 0 = does not demonstrate understanding). The learner’s answer is considered satisfactory if the prompt “demonstrates awareness of their own knowledge” This is to ensure the learner cannot proceed with the audio session simply by restating a keyword the model used back to it, as previous work showed that domain-specific words in reflections were not indicative of the quality of reflection, but rather reflective depth (Cloude et al., 2021).

We employed in-context learning (Dong et al., 2022) using examples in the prompt to guide the agent’s judgment. An example was if a learner is listening to content about Confucius, and they respond to a reflection prompt as “Confucius” to the model, this—while technically not incorrect, we guided learners to provide a more, detailed response to demonstrate synthesis of their knowledge to demonstrate a suitable reflection. For example, a response to a prompt with “Confucius’ teachings would be considered patriarchal by modern standards” demonstrates a learner’s understanding by combining a facet of what the content learned with how they contextualize the subject to present day.

It is important to note that the reflection prompts are different from quizzes. The learner does not need to mention everything they learned about the topic, only by demonstrating their ability to synthesize new understanding, the model deems their engagement and reflection on the material. The binary satisfactory/unsatisfactory classification acts as an elegant gate to guiding learner’s progress in real-time to capture reflective depth, as opposed to relying on keyword matching, while avoiding the complexity of multi-dimensional rubrics. In the Standard mode, our system continuously listens for interruptions but otherwise continues speaking until one occurs, and does not prompt the learner to reflect. During learner interactions, the system uses Deepgram for learner speech transcription, and Silero VAD ⁴ to detect when learners were speaking. Additionally, a fine-tuned SmolLM v2 model (Allal et al., 2025) predicts speech boundaries to support smooth turn-taking.

3.4 Podcast Interface

As shown in Figure 2, the web application displays a decorative wave, an abstract animated visualization of the generated speech that animates based on the volume and cadence of the AI podcaster’s

voice. When the podcaster is silent, the wave appears as a flat line of dots. Learners begin the session with their microphone automatically turned on after granting permission through their browser.

4 Methods

4.1 Sample

To build on the methods used by Do et al. (2025), we designed our study as an extension of their work and used the same source material and measurements. This study was approved by an Institutional Review Board and a total of 36 ($n=36$; 42% female) college students enrolled at universities in the United States were recruited through the Prolific online marketplace. Participants were pre-screened for English fluency and minimal prior knowledge of the subject (Introductory Philosophy). We also screened participants for technical requirements to ensure they had a working microphone and speaker for audio. One participant was excluded from our analysis due to adversarial responses to reflection prompts (e.g., “I hate bots and I hate them in the work place [sic]”). Due to the added length introduced by the interaction in the Reflection condition, we limited the scope to a single textbook, Introduction to Philosophy (Chapter 1) (Smith, 2022), and focused only on one subsection (Chapter 1.1), to ensure a manageable session duration while maintaining consistency with the original study design.

4.2 Procedure

The study consisted of a single 40-minute remote session. Before the session, participants were randomly assigned to the 1) Reflection or 2) Standard condition. Next, participants completed a brief demographic survey and were informed they would interact with an AI-generated podcast to learn about philosophy, after which we collected informed consent. Participants were then directed to a web application (described in Section 3) and guided through an interactive AI-generated podcast lasting approximately 15 minutes. Upon completion, the agent provided a verbal code and displayed a popup in the browser, enabling participants to proceed. They were redirected to a survey, where they completed the learning outcomes test and the User Experience Questionnaire (UEQ). Participants were compensated with \$10 USD after finishing the study.

⁴<https://github.com/snakers4/silero-vad>

4.3 Dependent Variables

4.3.1 Learning Outcomes

Learning was assessed using items from the post-chapter test bank from the OpenStax textbook⁵. The assessment items included seven multiple-choice questions from the Section 1.1 test bank. Due to the small number of multiple-choice questions for the single subsection, we included adapted three open-response items into multiple-choice questions, known as “Review Questions” (Smith, 2022), resulting in a total of 10 questions (see adapted items in Appendix B). Statistical analysis revealed no significant score differences between the original and supplemental items ($ps > .05$).

4.3.2 User Experience

User experience was measured using the User Experience Questionnaire (UEQ) (Laugwitz et al., 2008) immediately after the audio session, specifically the *Attractiveness* and *Stimulation* subscales (see Appendix A). These subscales gauge the overall appeal and engagement of user’s experience during the interaction, respectively. The UEQ employs a 7-point anchored Likert scale using adjective pairs such as “Annoying–Enjoyable” for *Attractiveness* and “Demotivating–Motivating” for *Stimulation*. We measured *Attractiveness* and *Stimulation* by averaging item scores within each subscale.

5 Results

Before conducting statistical analysis, we assessed whether our data adhered to a normal distribution using histograms and the D’Agostino-Pearson test, which suggested that our data were normally distributed across all variables ($K^2 = 2.56, p = .28$).

5.1 Research Question 1

To address our first research question, do interactive reflection prompts improve learning outcomes when incorporated into AI-generated podcasts compared to standard AI-generated podcasts, we calculated a two-sample t -test to compare whether there were differences in learning outcomes between Reflection and Standard conditions. The results suggested that there were no differences in learning outcomes between Reflection and Standard conditions, $t(34) = 0.89, p = 0.38, D = 0.29$.

⁵We do not provide the questions and answers for the knowledge retention questionnaires due to OpenStax policy. Verified educators from academic institutions may access test banks directly through OpenStax.

5.2 Research Question 2

To address our second research question, do interactive reflection prompts improve user experience when incorporated into AI-generated podcasts compared to standard AI-generated podcasts, we calculated 2 separate two-sample t -tests to compare whether there were differences in user experience subscales: attractiveness and stimulation. The results suggested there was a significant difference in *Attractiveness*, $t(34) = 2.26, p = 0.03, D = 0.75$, where the Standard condition rated the experience more favorably than the Reflection condition (Figure 3). Conversely, there were no significant differences in *Stimulation*, $t(34) = 1.31, p = 0.20, D = 0.44$, between the Standard and Reflection conditions. Descriptive statistics are in Table 1.

Table 1: Dependent variable descriptive statistics by condition.

	Learning	Attractiveness	Stimulation
Reflection	5.89 (1.94)	26.22 (4.58)	21.22 (3.68)
Standard	6.50 (2.06)	29.56 (4.00)	23.17 (4.87)

Note. Means and (standard deviations) are provided.

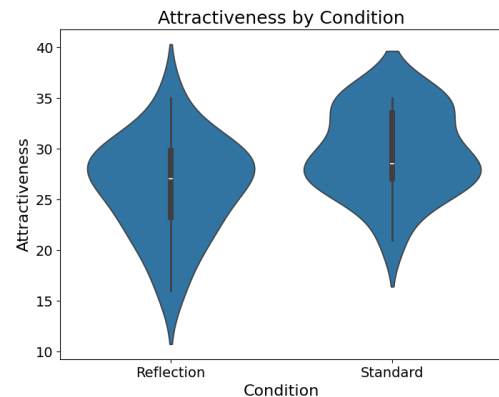


Figure 3: Attractiveness ratings across conditions.

6 Discussion

Personalized learning via interactions and reflection prompts are both recognized as valuable tools to enhance engagement and active learning (Sahronih et al., 2019) (Zhai et al., 2023). We implemented reflection prompts guided by (McAlpine et al., 1999)’s model of reflection and empirical literature suggesting that deeper reflection enhances learning, as supported by prior research that included no evaluation of responses in real-time (Cloude et al., 2021; Carpenter et al., 2021).

To build on this, our system applied a one-shot evaluation to judge learners' understanding based on their spoken reflections. Our study sought to explore whether integrating reflection prompts within an interactive, AI-generated educational podcast would improve user experience and learning outcomes compared to a standard, interactive AI-generated podcast.

Our first research question revealed that interactive, AI-generated podcasts with reflection prompts did not significantly improve learning outcomes compared to those without reflection prompts. This result contrasts with previous research, which found that reflection enhanced learning outcomes in game-based environments (Cloude et al., 2021; Carpenter et al., 2021). One explanation could be the nature of the prompt, which asked learners to recall factually relevant information rather than promoting reflection in relation to learning goals or planning, which encourages more thorough reflection and which is required by a game-based environment. Another possibility is that while reflection may be a useful tool for AI-generated podcasts, simply relying on the LLM to guide the reflection process without accounting for the learning goals or knowledge state of the learner may not be effective for promoting reflection with interactive AI-generated podcasts.

In our second research question, we found that the reflection prompts with an interactive, AI-generated podcast significantly reduced *Attractiveness* ratings compared to the standard AI-generated podcast condition. This indicated that the interactive elements in the Reflection condition may have disrupted the learners' flow and enjoyment during audio-based learning. This is different than previous research (Wang et al., 2025), which found that perceived interactivity and reflection boosted enjoyment and facilitated more active learning. The lack of significant differences for *Stimulation* suggests that the reflection intervention, despite its theoretical foundation in enhancing learning through reflective scaffolding (McAlpine et al., 1999), did not measurably improve user experience or learning outcomes in our study. Effective podcast-based reflections with LLMs likely require more detailed scaffolding, such as fine-tuning the model to provide automatic, tailored feedback that is based on learners' individual goals and current knowledge state. Future work should focus on developing stronger guidance methods to support reflection. This addresses a limitation identified by Do et al.,

who reported that participants desired "opportunities for active engagement" with AI-generated podcasts (2025), and suggests that while the specific implementation of reflection prompts may have detracted from the user experience, the general concept of interactive learning remains appealing to learners. The challenge appears to be finding the right balance between maintaining content flow and providing meaningful opportunities for reflection that effectively support learning.

6.1 Limitations

This study has important limitations to consider. First, our sample size of 36, which may limit the statistical power and generalizability of our results. Moreover, the focused scope of the content being taught (one section of a chapter) may not fully represent how reflection impacts learning across different subjects. Furthermore, our reflection responses were evaluated using a binary metric (understood/not understood) rather than evaluating the depth of reflection. This methodological constraint, though appropriate for our specific learning context, may have reduced the potential effectiveness of the reflection intervention compared to more elaborate implementations with different types of prompts. There are likely degrees to understanding which learning tools often fail to capture. Perhaps a human learner may feel more inclined to skip content they aren't understanding only to return to it later. Finally, the learner was exposed to the content for only 15 minutes, which may have reduced their learning and reflection due to the short nature of that task.

6.2 Future Work

Future research should explore alternative approaches for incorporating reflection and interaction in AI-generated podcasts. Developing adaptive reflection systems using LLMs that dynamically adjust based on learner engagement and metacognition would be a promising direction. Future work should investigate the use of LLMs for more fine-grained grading approaches for reflection quality, moving beyond binary assessments to evaluate responses with greater nuance. Additionally, investigating whether multi-modal data could better inform interaction and how to prompt reflections (e.g., eye movements, physiology, facial expressions, prior reflection quality, etc.) to enhance understanding.

References

- Loubna Ben Allal, Anton Lozhkov, Elie Bakouch, Gabriel Martín Blázquez, Guilherme Penedo, Lewis Tunstall, Andrés Marafioti, Hynek Kydlíček, Agustín Piqueres Lajarín, Vaibhav Srivastav, Joshua Lochner, Caleb Fahlgren, Xuan-Son Nguyen, Clémentine Fourier, Ben Burtenshaw, Hugo Larcher, Haojun Zhao, Cyril Zakka, Mathieu Morlon, and 3 others. 2025. [SmolLM2: When smol goes big – data-centric training of a small language model](#). *Preprint*, arXiv:2502.02737.
- Naomi Baron and Anne Mangen. 2021. [Doing the reading: The decline of long form reading in higher education](#). *Poetics Today*, 42:253–279.
- Dan Carpenter, Elizabeth Cloude, Jonathan Rowe, Roger Azevedo, and James Lester. 2021. [Investigating student reflection during game-based learning in middle grades science](#). In *LAK21: 11th International Learning Analytics and Knowledge Conference*, LAK21, page 280–291, New York, NY, USA. Association for Computing Machinery.
- Elizabeth Cloude, Dan Carpenter, Daryn A. Dever, James Lester, and Roger Azevedo. 2021. [Game-based learning analytics for supporting adolescents’ reflection](#). *Journal of Learning Analytics*, 8(2):51–72.
- Tiffany D. Do, Usama Bin Shafqat, Elise Ling, and Nikhil Sarda. 2025. [Paige: Examining learning outcomes and experiences with personalized ai-generated educational podcasts](#). In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (CHI)*, pages 1–10.
- Qingxiu Dong, Lei Li, Damai Dai, Ce Zheng, Jingyuan Ma, Rui Li, Heming Xia, Jingjing Xu, Zhiyong Wu, Tianyu Liu, and 1 others. 2022. [A survey on in-context learning](#). *arXiv preprint arXiv:2301.00234*.
- Russ d’Sa. 2024. [Openai and livekit partner to turn advanced voice into an api](#).
- Fangzhou Jin, Chin-Hsi Lin, and Chun Lai. 2025. [Modeling ai-assisted writing: How self-regulated learning influences writing outcomes](#). *Computers in Human Behavior*, 165:108538.
- Philippe Laban, Elicia Ye, Srujay Korlakunta, John Canny, and Marti Hearst. 2022. [Newspod: Automatic and interactive news podcasts](#). In *Proceedings of the 27th International Conference on Intelligent User Interfaces*, IUI ’22, page 691–706, New York, NY, USA. Association for Computing Machinery.
- Bettina Laugwitz, Theo Held, and Martin Schrepp. 2008. [Construction and evaluation of a user experience questionnaire](#). In *HCI and Usability for Education and Work: 4th Symposium of the Workgroup Human-Computer Interaction and Usability Engineering of the Austrian Computer Society, USAB 2008, Graz, Austria, November 20-21, 2008. Proceedings 4*, pages 63–76. Springer.
- Lynn McAlpine, Cynthia Weston, Catherine Beauchamp, C Wiseman, and Jacinthe Beauchamp. 1999. [Building a metacognitive model of reflection](#). *Higher education*, 37:105–131.
- Inge Molenaar, Susanne de Mooij, Roger Azevedo, Maria Bannert, Sanna Järvelä, and Dragan Gašević. 2023. [Measuring self-regulated learning and the role of ai: Five years of research using multimodal multichannel data](#). *Computers in Human Behavior*, 139:107540.
- Xuefei Ning, Zinan Lin, Zixuan Zhou, Zifu Wang, Huazhong Yang, and Yu Wang. 2023. [Skeleton-of-thought: Large language models can do parallel decoding](#). *Proceedings ENLSP-III*.
- Siti Sahronih, Agung Purwanto, and M. Syarif Sumantri. 2019. [The effect of interactive learning media on students’ science learning outcomes](#). In *Proceedings of the 2019 7th International Conference on Information and Education Technology*, ICIET 2019, page 20–24, New York, NY, USA. Association for Computing Machinery.
- Atikah Shemshack and Jonathan Michael Spector. 2020. [A systematic literature review of personalized learning terms](#). *Smart Learning Environments*, 7(1):33.
- Nathan Smith. 2022. [Introduction to Philosophy](#). OpenStax, Houston, Texas.
- Tanya Stivers, N. J. Enfield, Penelope Brown, Christina Englert, Makoto Hayashi, Trine Heinemann, Gertie Hoymann, Federico Rossano, Jan Peter de Ruiter, Kyung-Eun Yoon, and Stephen C. Levinson. 2009. [Universals and cultural variation in turn-taking in conversation](#). *Proceedings of the National Academy of Sciences*, 106(26):10587–10592.
- Feifei Wang, Alan C. K. Cheung, Ching Sing Chai, and Jin Liu. 2025. [Development and validation of the perceived interactivity of learner-ai interaction scale](#). *Education and Information Technologies*, 30(4):4607–4638.
- Yuchi Yahagi, Rintaro Chujo, Yuga Harada, Changyo Han, Kohei Sugiyama, and Takeshi Naemura. 2025. [Paperwave: Listening to research papers as conversational podcasts scripted by llm](#). In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems (CHI EA ’25)*, page 10, New York, NY, USA. ACM.
- Na Zhai, Yong Huang, Xiaomei Ma, and Jingchun Chen. 2023. [Can reflective interventions improve students’ academic achievement? a meta-analysis](#). *Thinking Skills and Creativity*, 49:101373.

A User Experience Questionnaire

All items were assessed on a 7-point scale, with the terms as anchors, adapted from (Laugwitz et al., 2008).

Attractiveness

Annoying – Enjoyable

Bad – Good

Unlikeable – Pleasing

Unpleasant – Pleasant

Unattractive – Attractive

Unfriendly – Friendly

Stimulation

Inferior – Valuable

Boring – Exciting

Not interesting – Interesting

Demotivating – Motivating

B Review Questions

We adapted three additional questions from the free-response items in (Smith, 2022) into multiple-choice format, alongside the chapter questions.

1. What characteristics are essential for being identified as a “sage”?

- a) Upholding social norms and exercising political power
- b) Seeking profound understanding through critical inquiry and providing foundational insights
- c) Mastering persuasive rhetoric and accumulating significant wealth
- d) Adhering to religious doctrines and conducting spiritual rituals

2. What does it mean for philosophy to “have an eye on the whole”?

- a) Rejection of traditional narratives through empirical investigation
- b) Fusion of mystical beliefs with systematic logical analysis
- c) Skeptical inquiry into established wisdom and foundational explanations of reality
- d) Emphasis on practical skills for societal and technological advancement

3. Which philosopher held that moral behavior and social harmony were linked to the natural order?

- a) Confucius
- b) Pythagoras
- c) Thales
- d) Yajnavalkya