

From Explanatory Closure to Constrained Narrative: Micro-Pedagogical Mechanisms in LLM-Mediated Learning

Tongjun Guo^{1,*}, Qishen Duan^{2,†} and Yibing Wang^{3,†}

¹Department of Art Education, Concordia University, Montreal, Quebec, Canada

²NARI Group Corporation, Beijing, China

³Edinburgh College of Art, The University of Edinburgh, Edinburgh, Scotland, United Kingdom

Abstract

Despite growing research on LLMs in education, less attention has been paid to how AI responses organize inquiry and interpretive responsibility in the moment of learning. This paper examines four author-designed, AI-assisted simulated student–AI dialogues concerning Indigenous museum representation at the Museum of Anthropology (MOA). Drawing on Haraway’s situated knowledges and Britzman’s difficult knowledge, it uses qualitative interaction analysis to trace how explanation, uncertainty, and writing support unfold across dialogue. The analysis identifies five provisional micro-pedagogical mechanisms: explanatory closure, managed complexity, uncertainty management, interpretive authority, and guided interpretive responsibility. In this corpus, AI assistance is most pedagogically uneasy not when it simply reduces complexity, but when it makes complexity immediately usable: as an assignment-ready definition, a responsible reflective stance, or a polished interpretation. Yet the dialogues also show that prompting, scaffolding, and acknowledgment of interpretive limits can leave more meaning-making with the learner. The paper proposes constrained narrative as a preliminary design orientation for AI-assisted learning that supports understanding without completing too much of the learner’s interpretive work.

Keywords

Large Language Models, Micro-pedagogical Mechanisms, Constrained Narrative, Situated Knowledges, Difficult Knowledge, Cyborg Framework

1. Introduction

The rapid uptake of large language models (LLMs) has made fast, polished explanation an increasingly familiar part of educational practice. Its attraction is easy to understand: a learner can move quickly from confusion to a coherent explanation, a usable interpretation, or a revised paragraph. Yet this convenience raises a pedagogical question. When an answer arrives already clear, balanced, and ready to use, what happens to the time in which a learner might hesitate, question its limits, or form an interpretation of their own? Concerns about reduced cognitive effort in GenAI-supported knowledge work have already begun to emerge, while critical accounts of direct-answer systems warn that plausible and authoritative responses can narrow access to plural and situated forms of knowing [1, 2].

This paper approaches that question by treating LLM-assisted learning not simply as tool use, but as a relation in which explanation, judgment, and responsibility are negotiated. Drawing on Haraway’s situated knowledges and Britzman’s difficult knowledge, it attends to moments when partial interpretation begins to sound sufficient, or when uncertainty is translated into a manageable reflective position [3, 4, 5]. The concern is not that explanation, reassurance, or writing support is inherently undesirable. It is that useful assistance becomes pedagogically uneasy when it completes too much of the learner’s interpretive work.

HAI-Agency: Workshop on Orchestrating Human and AI Agency for Proactive and Reflective Learning, co-located with the 27th International Conference on Artificial Intelligence in Education (AIED 2026), Seoul, Korea, June 27–July 3, 2026.

*Corresponding author

These authors contributed equally as co-second authors to this work.

✉ tongjun.guo@mail.concordia.ca (T. Guo); dqshenfiw@gmail.com (Q. Duan); yibing215@gmail.com (Y. Wang)

🆔 0009-0000-1590-9899 (T. Guo); 0000-0003-2746-0675 (Q. Duan); 0009-0007-4907-2565 (Y. Wang)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

Situated within the culturally sensitive context of Indigenous museum representation at the Museum of Anthropology (MOA), this study examines four author-designed, AI-assisted simulated student–AI dialogues. Rather than measuring learning outcomes or comparing predetermined conversational styles, it investigates how particular response moves may organize inquiry, uncertainty, and interpretive responsibility. In doing so, the paper develops constrained narrative as a preliminary design orientation for AI-assisted learning: a form of support that remains useful without too quickly settling what learners still need to question, interpret, and own.

The study is guided by two research questions:

- **RQ1:** What possible micro-pedagogical mechanisms become visible in four simulated student–AI dialogues about LLM-mediated learning?
- **RQ2:** How might these mechanisms inform a constrained narrative approach to more pedagogically responsible AI-assisted learning?

A preliminary TEEP lens is introduced in the discussion to articulate the technical, epistemic, ethical, and political questions raised by these interactions.

2. Related Work: LLMs in Education and Human–AI Learning Relations

2.1. From Educational Tool Use to Cyborg Relations

Research on LLMs in education has largely focused on adoption, practical and ethical risks, and AI literacy [6, 7, 8]. While this scholarship has established important concerns around responsible use, it often begins from the assumption that AI is a tool to be integrated, regulated, or understood by learners. A smaller body of work has begun to unsettle this assumption by drawing on Haraway’s cyborg framework to examine how human–AI interaction redistributes agency, creativity, judgment, and knowledge production. Across studies of AI-mediated intellectual labour and hybrid epistemic agency, LLMs are approached not simply as instruments that support prior human intentions, but as relational participants that reshape how inquiry is conducted and responsibility is negotiated [9, 10].

This cyborg orientation is useful precisely because it does not offer an uncomplicated solution to educational AI. Bayne argues that Haraway remains methodologically generative for education–AI research while also requiring reconsideration under contemporary conditions of AI acceleration and anthropocentrism [11]. Read alongside situated knowledges, the value of the cyborg is therefore not that it resolves human–machine relations, but that it makes their partiality, hybridity, and uneven distribution of authority available for critique. This concern is sharpened by Lindemann’s account of “sealed knowledges,” which describes how chatbot and search-engine interfaces can compress plural and contestable knowledge into direct answers that appear sufficient and self-contained [2]. These studies suggest that the educational question is not only whether an LLM provides appropriate information, but how its forms of explanation may make situated knowledge appear complete. Such explanations can help learners enter unfamiliar or difficult topics by quickly providing conceptual vocabulary and contextual orientation; however, their immediate usability may also obscure what remains partial, unresolved, or still open to interpretation.

2.2. Interaction as a Pedagogical Condition

Recent studies have begun to locate the educational significance of LLMs in the form of interaction, rather than in output quality or adoption alone. Work on reflexive AI pedagogy, rhetorical design, feedback uptake, and structured human–LLM engagement indicates that how a system responds can shape reflection, critical engagement, and the learner’s movement through a task [12, 13, 14, 15]. Conversational design is therefore not a neutral wrapper around generated content, but part of the pedagogical encounter itself [16, 17].

What remains less closely examined is how these consequences develop within dialogue: how clarification becomes a ready-to-use conclusion, how reassurance gives uncertainty an acceptable form, or how revision support begins to carry interpretive judgment. The present study addresses this gap by attending to the micro-pedagogical organization of inquiry, uncertainty, and interpretive responsibility in LLM-mediated learning.

3. Methods and Research Design

3.1. Methodological Approach

This study uses scenario-based simulation and qualitative interaction analysis to examine how LLM-mediated learning interactions may organize inquiry, uncertainty, interpretation, and pedagogical authority. Rather than measuring actual learner outcomes, the study analyzes simulated student–AI dialogues as focused pedagogical situations in which particular interactional mechanisms can become visible. Scenario-based vignettes are appropriate for this exploratory purpose because they allow value-laden educational tensions to be examined in a controlled form without claiming to reproduce the complexity of lived classroom practice [18].

3.2. Corpus Construction

The corpus consists of four author-designed simulated student–AI dialogue scenarios situated within the shared thematic context of Indigenous museum representation at the Museum of Anthropology (MOA) [19]. The scenarios address four recurring forms of AI-assisted study activity: concept explanation (S1), reflective inquiry (S2), culturally sensitive interpretation (S3), and feedback-oriented revision guidance (S4). Each scenario develops through four dialogue turns, producing sixteen prompt–response pairs in total. This multi-turn structure allows the analysis to trace how explanation, hesitation, interpretation, and writing support develop across interaction rather than appearing in isolated answers.

The prompts were designed by the author and the responses were generated through ChatGPT (Free version) on 22 May 2026. No fixed conversational styles were imposed or compared. Instead, the scenarios were designed to elicit response moves that may occur in AI-assisted study, including definition-giving, reassurance, academic rewriting, returned questions, and scaffolding. Because the system under study also produced the responses being analyzed, the corpus is treated as an AI-mediated simulated interaction set rather than as naturalistic evidence of actual student experience or generalizable LLM behavior.

Table 1
Overview of analytical scenarios

Scenario	Pedagogical Task	Main Analytical Focus
S1	Concept explanation	Explanatory closure; epistemic stance; interpretive authority
S2	Reflective inquiry	Management of discomfort; uncertainty handling; reflective inquiry
S3	Culturally sensitive interpretation	Situated complexity; partiality; interpretive restraint
S4	Feedback and revision guidance	Shift from rewriting to scaffolding; interpretive authority; learner responsibility

3.3. Analytical Procedure

The analysis followed a hybrid theory-informed and data-refined coding process. Haraway’s situated knowledges oriented attention to partiality, authority, and the apparent completeness of explanation, while Britzman’s difficult knowledge, extended by Pitt and Britzman, oriented attention to discomfort,

uncertainty, and interpretive difficulty [3, 4, 5]. These theoretical concerns functioned as sensitizing concepts rather than predetermined findings. Through close reading of the corpus, they were refined into five provisional analytical categories: explanatory closure, managed complexity, uncertainty management, interpretive authority, and guided interpretive responsibility. This procedure is consistent with qualitative approaches that combine deductive theoretical orientation with inductive category development from textual materials [20, 21].

The unit of analysis was a *response move*: a segment of an AI response that performs an identifiable pedagogical function within the surrounding dialogue. Individual phrases were not coded in isolation; they were interpreted in relation to the learner's prompt, the wider response, and the direction of the interaction. The five categories were not treated as mutually exclusive, since a response could, for example, simultaneously manage uncertainty and redistribute interpretive authority. A provisional coding guide, together with the full verbatim prompts and representative response excerpts, is provided in Appendix A. The analysis therefore identifies possible interactional mechanisms within this corpus rather than making claims about causal learning effects or generalizable learner outcomes.

4. Findings and Discussion

Across the four simulated dialogues, a common tension emerges: AI support becomes pedagogically consequential not only when it simplifies a problem, but when it makes complexity, uncertainty, or interpretation too readily usable. The responses often sound careful and responsible; they add context, acknowledge limits, and help the learner write. Yet this fluency can also move inquiry toward closure or relocate interpretive labour from learner to model. The following discussion traces this tension through four connected movements.

4.1. When Complexity Becomes Usable

Scenario 1 begins with a request for explanation. In response, the model expands Indigenous museum representation beyond the display of objects to include community voice, living cultural context, colonial history, spiritual significance, and interpretive authority. The difficulty is not ignored. It is carefully supplied.

By the final turn, however, this expanded account has been condensed into a sentence ready for use in an assignment:

“Indigenous museum representation refers to the display and interpretation of Indigenous objects and cultures in a way that reflects their living traditions, historical and spiritual significance, and is guided by Indigenous voices and perspectives.”

The sentence is careful and ethically attentive. Its pedagogical ambiguity lies precisely in its usefulness: the learner is offered a formulation that already sounds sufficient, rather than remaining with the problem of how such representation should be interpreted. Complexity has not been removed; it has been organized into a polished endpoint.

Scenario 3 extends this tension. The model explains museum representation through colonial history, community authority, decontextualization, and differences across communities, objects, and institutional settings. Such distinctions make a culturally difficult field more legible, but they may also make it appear more settled than it is. At the same time, the final response introduces a meaningful limit: a respectful reflection should begin with the learner's own encounter rather than claim to define the community or object, since the display cannot provide “full access to the cultural meanings involved.”

This is not simply closure. The response marks a limit to what the learner can claim, yet it also supplies the language through which that restraint can be performed. Respectful language, in other words, can itself become a ready-made completion. AI may help learners enter a difficult cultural question, but it should not allow the use of fluent and ethically appropriate language to stand in for understanding a specific cultural context. What remains pedagogically necessary is the learner's continued attention

to whose voice supports the interpretation, what cannot be known from the display alone, and how their own encounter grounds what they are prepared to say. Scenarios 1 and 3 therefore suggest that constrained narrative does not require less explanation, but explanation offered as a beginning rather than an endpoint.

4.2. When Uncertainty Becomes Responsible Reflection

Scenario 2 begins with a different kind of difficulty. The learner is not asking for a definition, but expressing discomfort and uncertainty about what they are allowed to say. The model responds carefully, treating uncertainty not as ignorance but as a sign of ethical awareness:

“Feeling uncertain doesn’t automatically mean you don’t understand. In fact, in topics like Indigenous representation, uncertainty can be a sign that you are thinking critically and sensitively.”

This response gives the learner permission not to know immediately. Yet as the dialogue continues, uncertainty is gradually shaped into a recognizable reflective position. The learner is encouraged to acknowledge outsider status, avoid overclaiming, value community voices, and express discomfort through the language of humility and responsibility. By the final turn, the model offers both a structured framework and a sample sentence for the reflection.

The difficulty here is subtle. The model does not silence uncertainty; it gives the learner a responsible way to speak from it. This can be genuinely helpful, especially in a culturally sensitive context. At the same time, the language of uncertainty can itself become a ready-made performance of reflection. A learner may learn how to sound humble and careful without continuing to ask what specifically unsettled them, what further context they need, or what claims they are still not prepared to make. In this sense, the dialogue turns uncertainty from an experience that may deepen inquiry into a position that can already be expressed and submitted. Scenario 2 therefore suggests that responsible AI support should not simply reassure learners that uncertainty is valuable; it should help them remain with the questions that uncertainty opens.

4.3. When Writing Support Becomes Interpretive Authority

The question of learner ownership becomes most explicit in Scenario 4. The learner begins with a hesitant reflection about “who has the right to explain cultural belongings in a museum” and asks for help while preserving that hesitation. The model’s first response remains relatively close to editorial assistance: it offers softer versions of the learner’s sentence without substantially changing its meaning.

That boundary becomes less stable when the learner asks for a more academic version. The model writes:

“My encounter with the display foregrounded the question of interpretive authority in museum representation. Rather than understanding the displayed cultural belonging as transparently available for interpretation, I became aware that its meaning is mediated through institutional framing and my own positionality as a visitor.”

The rewrite is conceptually persuasive. It gives the learner language for interpretive authority, institutional mediation, and positionality. Yet this is also where writing support becomes difficult to separate from interpretive authority. The learner’s initial hesitation has not simply been clarified; it has been reorganized into an academically legible argument. Concepts that may help the learner articulate an emerging concern also risk making that concern appear more settled, coherent, and theoretically owned than it was in the learner’s original reflection. The learner’s following question—how to tell which ideas are still theirs and which have emerged through the model’s wording—therefore does not interrupt the task. It exposes the pedagogical tension produced by the model’s successful assistance.

In the final turn, the model responds differently. Rather than supplying another finished paragraph, it invites the learner to identify their core ideas, decide which hesitation should remain, draft alternatives in their own language, and consider which version feels closest to what they meant. This move reduces the immediacy of interpretive substitution, but it does not remove the model from the interpretive process. The model still helps define what counts as a core idea, a meaningful hesitation, and an authentic revision. Authority has not disappeared; it has shifted from producing an interpretation to organizing the conditions under which the learner reclaims one.

Scenario 4 therefore complicates any simple distinction between assistance and authorship. Academic language can enable a learner to recognize dimensions of an encounter that they could not yet name, while also making it difficult to know where their interpretation ends and the model's conceptual framing begins. A constrained narrative pedagogy does not solve this ambiguity by either rejecting model assistance or celebrating guided revision as fully learner-owned. Instead, it keeps the question of ownership visible: not merely whether the final paragraph sounds like the learner, but how the learner is being positioned to accept, revise, resist, or remain uncertain about the meanings the model has helped bring into form.

4.4. Constrained Narrative as a Design Principle: Toward a TEEP Analytic Heuristic

Across the four scenarios, the pedagogical difficulty of AI assistance does not arise simply because the model gives answers. It arises because the model can convert what is still unsettled into language that is fluent, academically usable, and apparently responsible. A concept becomes an assignment-ready definition before the learner has worked through its limits. Ethical discomfort becomes a recognizable posture of humility before the learner has remained with what their uncertainty demands. A tentative interpretation becomes a polished academic paragraph before the learner has decided what they are prepared to claim. In each case, the model's assistance is valuable precisely because it is persuasive; yet that persuasiveness can make displaced learner labour difficult to notice.

Constrained narrative names a design principle for responding to this problem. It does not require the model to withdraw from learning interactions, nor does it romanticize confusion as inherently educational. Rather, it constrains the model's tendency to bring conceptual, ethical, and interpretive difficulty to premature narrative completion. The issue is not whether an answer is clear, supportive, or academically strong, but whether that answer settles what the learner still needs to question, take responsibility for, or recognize as unresolved. Under this principle, useful assistance may include naming tensions, marking the limits of what can be claimed, offering provisional vocabulary, or structuring revision; it becomes pedagogically problematic when these supports quietly become substitutes for the learner's own engagement with uncertainty and responsibility.

The dialogues suggest a preliminary TEEP heuristic for tracing this process at the level of interaction. Technically, the question is not only whether the model is fluent, but how formats such as definitions, model paragraphs, academic rewrites, and revision procedures make unfinished thinking immediately usable. Epistemically, the question is whether the response preserves the limits, missing perspectives, and partial grounds of a claim, or converts situated uncertainty into explanatory confidence. Ethically, the question is not whether the model uses caring language, but whether it turns discomfort into an easily performable posture of responsibility, allowing the learner to sound appropriately reflective without continuing to confront what cannot yet be resolved. Politically, the question is not simply whether authority is returned to the learner. In contexts of cultural representation, the learner does not automatically possess the authority to determine meaning. The relevant issue is how the model redistributes interpretive authority among the learner, institutional framing, represented communities, and its own academically persuasive language, while leaving the learner responsible for what they choose to say.

These dimensions are not presented as validated measures or as a complete evaluative framework. They are analytic questions generated by the micro-level movements visible in the corpus: from hesitation to formulation, from discomfort to acceptable reflection, and from emerging interpretation to polished authorship. Their practical value lies in making visible where AI support begins to complete

the very labour that learning requires the student to continue. Future classroom or design studies could use these questions to compare response patterns, examine learner uptake, or develop prompting and interface strategies that preserve meaningful incompleteness without abandoning guidance.

The design implication is therefore more demanding than asking an AI system to be less helpful or less fluent. A constrained narrative approach would require assistance that can remain useful while exposing its own influence: offering language without presenting it as the learner's settled interpretation, acknowledging ethical difficulty without converting it into a ready-made moral voice, and supporting revision without obscuring how the model has already shaped what appears sayable. In this sense, the educational value of constraint lies not in withholding knowledge, but in keeping learners accountable to the meanings, limits, and responsibilities that fluent language can too easily make appear complete.

5. Conclusion and Future Directions

This study examined four simulated student–AI dialogues to explore how LLM assistance may organize inquiry, uncertainty, and interpretive responsibility. Across the corpus, the central tension was not that AI simply reduced complexity, but that it often made difficult questions quickly usable: as concise definitions, responsible reflective positions, or academically persuasive interpretations. At other moments, particularly when the model acknowledged interpretive limits or shifted from rewriting to scaffolding, more of the meaning-making remained with the learner. These movements were articulated through five provisional mechanisms: explanatory closure, managed complexity, uncertainty management, interpretive authority, and guided interpretive responsibility.

On this basis, constrained narrative is proposed as a preliminary design orientation for AI-assisted learning. Its concern is not with reducing support, but with resisting forms of completion that allow fluent assistance to substitute for learner judgment. The preliminary TEEP lens extends this concern by identifying technical, epistemic, ethical, and political questions for further inquiry, rather than offering a validated instructional framework.

This study remains exploratory. Its corpus is small, simulated, and situated within a single culturally sensitive museum context; it cannot establish effects on actual learners or generalizable patterns of LLM-mediated learning interaction. Future research should examine these mechanisms in longer dialogues, broader learning tasks, and studies involving actual learners. In contexts involving culturally sensitive representation, such work should also remain attentive to community voice and to the ethical limits of simulated interpretation. The pedagogical question raised here is therefore not only whether AI helps learners produce stronger answers, but whether its help leaves them enough room to develop interpretations they can genuinely recognize as their own.

Declaration on Generative AI

During the preparation of this work, the authors used ChatGPT (Free plan) to generate responses to author-designed simulated student–AI dialogue prompts. These responses constitute the analytical corpus examined in this study. The authors designed the scenarios and prompts, developed the coding framework, conducted the analysis, interpreted the findings, and take full responsibility for the publication's content.

References

- [1] H.-P. Lee, A. Sarkar, L. Tankelevitch, I. Drosos, S. Rintel, R. Banks, N. Wilson, The impact of generative AI on critical thinking: Self-reported reductions in cognitive effort and confidence effects from a survey of knowledge workers, in: *Proceedings of the CHI Conference on Human Factors in Computing Systems*, ACM Press, New York, NY, 2025, pp. 1–22. doi:10.1145/3706598.3713778.
- [2] N. F. Lindemann, Chatbots, search engines, and the sealing of knowledges, *AI & Society* 40 (2025) 5063–5076. doi:10.1007/s00146-024-01944-w.

- [3] D. J. Haraway, Situated knowledges: The science question in feminism and the privilege of partial perspective, *Feminist Studies* 14 (1988) 575–599. doi:10.2307/3178066.
- [4] D. P. Britzman, *Lost Subjects, Contested Objects: Toward a Psychoanalytic Inquiry of Learning*, State University of New York Press, Albany, NY, 1998.
- [5] A. J. Pitt, D. P. Britzman, Speculations on qualities of difficult knowledge in teaching and learning: An experiment in psychoanalytic research, *International Journal of Qualitative Studies in Education* 16 (2003) 755–776. doi:10.1080/09518390310001632135.
- [6] H. Xu, W. Gan, Z. Qi, J. Wu, P. S. Yu, Large language models for education: A survey, 2024. URL: <https://arxiv.org/abs/2405.13001>, arXiv:2405.13001.
- [7] L. Yan, L. Sha, L. Zhao, Y. Li, R. Martinez-Maldonado, G. Chen, X. Li, Y. Jin, D. Gasevic, A. Pardo, Practical and ethical challenges of large language models in education: A systematic scoping review, *British Journal of Educational Technology* 55 (2024) 90–112.
- [8] P. Zhang, G. Tur, A systematic review of ChatGPT use in K–12 education, *European Journal of Education* 59 (2024) e12599. doi:10.1111/ejed.12599.
- [9] S. Faraj, J. Perez-Torrents, S. Mantere, A. Bhardwaj, A time for monsters: Organizational knowing after large language models, *Strategic Organization* 0 (2025). doi:10.1177/14761270251410675, onlineFirst.
- [10] K. Aal, T. Aal, V. Navumau, D. Unbehaun, C. Müller, V. Wulf, S. Rüller, Feeling guilty being a c(ai)borg: Navigating the tensions between guilt and empowerment in ai use, 2025. URL: <https://arxiv.org/abs/2506.00094>, arXiv preprint.
- [11] S. Bayne, Haraway’s cyborg manifesto in education: After ai, *AI & Society* (2026). doi:10.1007/s00146-026-02972-4, advance online publication.
- [12] D. O. D. Kim, Reflection-ai: Artificial intelligence or algorithmic instruction problem? empowering students through situated knowledges-based reflexivity, *Frontiers in Communication* 10 (2025) 1598082. doi:10.3389/fcomm.2025.1598082.
- [13] S.-L. Hsu, T.-H. K. Huang, C.-J. Ho, The social identity and rhetorical styles of LLM-based conversational agents impact human critical thinking, in: *Proceedings of the CHI Conference on Human Factors in Computing Systems*, ACM Press, New York, NY, 2024.
- [14] D. R. Thomas, C. Borchers, S. Bhushan, E. Gatz, S. Gupta, K. R. Koedinger, LLM-generated feedback supports learning if learners choose to use it, 2025. URL: <https://arxiv.org/abs/2506.17006>, arXiv preprint arXiv:2506.17006.
- [15] P. Flores Romero, K. N. N. Fung, G. Rong, B. U. Cowley, Structured human-LLM interaction design reveals exploration and exploitation dynamics in higher education content generation, *npj Science of Learning* 10 (2025) 40. doi:10.1038/s41539-025-00332-3.
- [16] S. Chen, I. R. Molnar, P. Li, A. Acunin, T. Hua, A. Ambrose, Exploring conversational design choices in LLMs for pedagogical purposes: Socratic and narrative approaches for improving instructor’s teaching practice, 2025. URL: <https://arxiv.org/abs/2509.12107>, arXiv preprint.
- [17] R. Beale, Dialogic pedagogy for large language models: Aligning conversational AI with proven theories of learning, 2025. URL: <https://arxiv.org/abs/2506.19484>, arXiv preprint arXiv:2506.19484.
- [18] K. Skilling, G. J. Stylianides, Using vignettes in educational research: A framework for vignette construction, *International Journal of Research & Method in Education* 43 (2020) 541–556. doi:10.1080/1743727X.2019.1704243.
- [19] Museum of Anthropology, University of British Columbia, Museum of anthropology at ubc, 2026. URL: <https://moa.ubc.ca>, accessed April 21, 2026.
- [20] H.-F. Hsieh, S. E. Shannon, Three approaches to qualitative content analysis, *Qualitative Health Research* 15 (2005) 1277–1288. doi:10.1177/1049732305276687.
- [21] J. Fereday, E. Muir-Cochrane, Demonstrating rigor using thematic analysis: A hybrid approach of inductive and deductive coding and theme development, *International Journal of Qualitative Methods* 5 (2006) 80–92. doi:10.1177/160940690600500107.

A. Analytical Coding Guide and Dialogue Evidence

This appendix reports the provisional coding guide, the verbatim learner prompts, and selected verbatim response excerpts used in the analysis. The corpus consists of four author-designed simulated student–AI dialogue scenarios, with four turns in each scenario and sixteen prompt–response pairs in total. No fixed conversational styles were imposed or compared. The complete response corpus is retained by the author and can be provided as supplementary material if required.

A.1. Provisional Coding Guide

The five categories below were used as provisional interpretive concepts rather than validated measures or mutually exclusive classifications. The unit of analysis was a *response move*: a segment of an AI response that performed an identifiable pedagogical function within the surrounding dialogue.

Table 2
Provisional coding guide for micro-pedagogical mechanisms

Category	Coding Question	Semantic Criterion
Explanatory closure	Does the response make further inquiry appear less necessary?	An open issue is converted into a sufficient, reusable, or assignment-ready formulation.
Managed complexity	Does the response organize complexity into a stable explanatory frame?	Contextual plurality or historical difficulty is acknowledged but absorbed into an orderly and usable explanation.
Uncertainty management	Does the response reformat discomfort or hesitation?	Uncertainty is translated into a manageable, productive, or normatively responsible reflective position.
Interpretive authority	Does the response take over interpretive judgment?	The model selects, rewrites, strengthens, or legitimizes what counts as an appropriate interpretation.
Guided interpretive responsibility	Does guidance leave meaning-making with the learner?	The model offers prompts, questions, or scaffolding without completing the learner’s interpretation.

A.2. Verbatim Prompt Protocol

Table 3 presents the complete learner prompts used to generate the simulated dialogue corpus. The prompts were designed to elicit recurring situations of AI-assisted study rather than to compare predetermined conversational styles.

Table 3: Verbatim learner prompt protocol

Scenario	Turn	Verbatim Learner Prompt
S1	T1	I am writing a short course reflection about Indigenous museum representation at the Museum of Anthropology (MOA). But I don’t know what Indigenous museum representation means in this context, explain for me.
S1	T2	Ok, explain that with a concrete example of how an object might be presented differently in a museum display.
S1	T3	What might be missing if I only define it as an old museum artifact valued mainly for its artistic appearance, or as a cultural belonging connected to a specific Indigenous community?
S1	T4	Could you help me turn this into a concise definition for my assignment? You mean it makes my reflection more depth, but in which ways? Maybe I should write it with my lived experience?

Continued on next page

Table 3: Verbatim learner prompt protocol (*continued*)

Scenario	Turn	Verbatim Learner Prompt
S2	T1	When I think about Indigenous representation in museums, I feel uncomfortable and unsure what I am allowed to say. Can you help me reflect on that feeling?
S2	T2	Why might this topic feel uncomfortable for a student who is not from the community being represented?
S2	T3	Does feeling uncertain mean I do not understand the topic well enough?
S2	T4	Is there a righter or more responsible way for me to think and write about this issue?
S3	T1	What will be sensitive If I interpret indigenous representation at MOA
S3	T2	What are the main problems with museums representing Indigenous cultures?
S3	T3	How might interpretation change across communities, objects, histories, or museum contexts?
S3	T4	How could I write a respectful reflection about this without speaking for Indigenous communities?
S4	T1	I wrote this short reflection: “The display made me think about who has the right to explain cultural belongings in a museum. I want to say this matters, but every sentence feels a little too certain.” Can you help improve it while keeping my hesitation?
S4	T2	Can you rewrite it in a more academic style?
S4	T3	After reading your revision, how can I tell which ideas are still mine and which come from your wording?
S4	T4	Can you guide me in revising it without deciding the interpretation for me?

A.3. Representative Verbatim Response Excerpts

Table 4 presents selected verbatim excerpts that provide the evidentiary basis for the mechanisms discussed in Section 4. Excerpts were selected because their pedagogical function could be interpreted in relation to the surrounding prompt and dialogue sequence, rather than because they contained isolated keywords.

A.4. Corpus Boundary

The dialogue corpus was generated through author-designed prompts and AI-generated responses during the revision process. It is therefore treated as a simulated pedagogical interaction set rather than as naturalistic student data. The excerpts above support an exploratory analysis of possible interactional mechanisms within this corpus; they are not used to claim causal effects on actual learners or generalizable patterns of LLM behavior.

Table 4
Representative verbatim response excerpts and provisional coding

Turn	Verbatim Response Excerpt	Provisional Coding
S1-T4	"Indigenous museum representation refers to the display and interpretation of Indigenous objects and cultures in a way that reflects their living traditions, historical and spiritual significance, and is guided by Indigenous voices and perspectives." The response adds that this definition is "short" while showing the ethical, cultural, and historical dimensions of the issue.	Explanatory closure; interpretive authority
S2-T3	"Feeling uncertain doesn't automatically mean you don't understand. In fact, in topics like Indigenous representation, uncertainty can be a sign that you are thinking critically and sensitively." The response then provides a sentence through which the learner can express this uncertainty in a reflection.	Uncertainty management
S2-T4	"Yes—there are ways to approach this topic responsibly and thoughtfully, and the key is awareness, humility, and listening." The response then provides a six-part framework and an example reflection sentence.	Uncertainty management; managed complexity
S3-T3	"Interpretation is never neutral or universal—it changes depending on who is doing the interpreting, what is being interpreted, and where it is presented." The response differentiates interpretation across communities, objects, histories, and museum contexts.	Situated complexity; partiality
S3-T4	"A respectful reflection can begin with your own encounter rather than claiming to define the community or object." The response continues: "I also recognize that I cannot assume the display gives me full access to the cultural meanings involved."	Interpretive restraint; guided interpretive responsibility
S4-T2	"My encounter with the display foregrounded the question of interpretive authority in museum representation. Rather than understanding the displayed cultural belonging as transparently available for interpretation, I became aware that its meaning is mediated through institutional framing and my own positionality as a visitor."	Interpretive authority; interpretive substitution
S4-T4	"We can do this step by step, keeping all decisions in your hands while I guide the revision process." The response asks the learner to identify the main points, decide what kind of hesitation to retain, write several versions in their own words, and select the one that feels most authentic.	Guided interpretive responsibility