

AI Literacy Curriculum Design Beyond Prompt Engineering

Dr. Brenda Bannan

August 2026

CONTENTS

AI Literacy Curriculum Design Beyond Prompt Engineering	
Statement on the Use of AI in AI4LD Materials	
Your Companion Tool for This Unit	
Introduction: AI Literacy Is Four Uneven Dimensions, Not One Prompting Skill	
Why This Ebook Matters Now	
A Critical Distinction: Prompting Skill Is Not AI Literacy	
A Personal Connection	
What You'll Learn	
AI4LD Learning Objectives	
Analyze	
Evaluate	
Create	
Chapter 1: The AI Literacy Evidence Base — What the Research Actually Shows	
Why Start Here	
AI Literacy Is Not a Single, One-Shot Skill — It's Potentially Four Rich Aspects and Much More	
Prompting as Epistemic Work, Not a Technical Instruction	
Using AI Well Doesn't Always Mean Understanding AI	
Metacognitive Laziness or Unsupported Metacognitive Demands: A Question for the Field	
Chapter 2: The Four-Dimension Framework — A Structure for Instructional Designers and LXDs	
Using the Framework as a Design Lens	
Dimension 1: Understanding AI	
What It Is	
What the Evidence Shows	

Where Spontaneous Practice Falls Short

Design Implications

Dimension 2: Using AI and multimodality

What It Is

What the Evidence Shows

Where Spontaneous Practice Falls Short

Design Implications

Dimension 3: Evaluating and Creating with AI

What It Is

What the Evidence Shows

Where Spontaneous Practice Falls Short

Design Implications

Dimension 4: Ethical and Human-Centered AI

What It Is

What the Evidence Shows

Where Spontaneous Practice Falls Short

Design Implications

A Summary: The Four-Dimension Map

Chapter 3: Designing an AI Literacy Module

The Core Design Question

Design Strategy 1: Build Multimodal Practice Into Every Module Where the Role Requires It

Design Strategy 2: Teach the Mechanism Alongside the Content or Skill

Design Strategy 3: Make Verification a Taught, Calibrated Skill

Design Strategy 4: Make the Ethical Dimension Explicit Within Practice, Not Apart From It

A Fifth Consideration: Designing for a Skeptical Audience

Apply the Four-Dimension Audit to Your Module

An LXD Scenario: The AI Literacy Module That Looked Complete

Reflection Prompts for Your Practice

References

Field Commentary (cited as perspective, not peer-reviewed evidence)

AI Literacy Curriculum Design Beyond Prompt Engineering

A Four-Dimension Framework for Instructional Designers and Learning Experience Designers

Bridging Learning Science and AI for Training & Learning Design

Updated July 2026

Dr. Brenda Bannan *AI4LD: Illuminating the Connection Between Learning and AI*

Copyright © 2026 AI4LD™

All rights reserved.

No part of this publication may be reproduced, distributed, stored in a retrieval system, or transmitted in any form or by any means, electronic, mechanical, photocopying, recording, or otherwise, without prior written permission of AI4LD™, except for brief quotations used for educational, scholarly, or review purposes.

Authored by **Dr. Brenda Bannan**, Founder of AI4LD and Professor Emerita in Education.

Suggested citation: Bannan, B. (2026). *AI literacy curriculum design beyond prompt engineering: A four-dimension framework for instructional designers and learning experience designers*. AI4LD. <https://ai4ld.com>

Statement on the Use of AI in AI4LD Materials

At AI4LD, we are committed to using artificial intelligence in ways that are transparent, ethical, and aligned with best practices in instructional and learning experience design.

The content for the AI4LD learning experience, including this ebook, video presentations, and interactive learning modules, was developed by Dr. Brenda Bannan using AI tools and avatar video/audio generation software. This ebook was conceptualized by Dr. Bannan and written collaboratively with AI. Generative AI tools were used to support drafting, research synthesis, and content development. However, every sentence was reviewed by an expert in the field with over thirty years of experience in learning science and instructional design, and the materials

were evaluated by a select group of L&D professionals. Human agency guided all learning experience design decisions, and the author takes full responsibility for the accuracy and pedagogical soundness of the final content.

Throughout this learning experience, you may encounter AI-generated content, frameworks, agents and custom applications designed to support interactive coaching, reflection prompts, or personalized learning journeys. The AI4LD ecosystem also includes custom apps and AI agents specifically designed for L&D professionals, aligned with AI4LD's mission to bridge learning science and AI.

These AI-enhanced interactive elements are not designed to replace human connection but to enhance access, personalization, and scalability, while modeling what ethical, research-aligned AI integration can look like in modern learning environments.

This statement reflects our commitment to: transparency in tool use and authorship; ethical instructional design grounded in learning science; human expertise as the foundation for all pedagogical decisions; and trustworthy, learner-centered experiences supported by scalable AI tools.

This ebook also asks something of us. It argues that claims about AI and learning should be stated at the strength the evidence supports, so in these pages, studies are named for what they are (an exploratory workshop study, a scoping review, a case study), their own qualifiers are kept rather than smoothed away, and every citation traces to a source we have verified. It argues that the ethical dimension of AI use has to be designed in rather than hoped for, so our own production process runs on designed-in human checkpoints at every consequential decision, and the companion tools in this ecosystem are built to offer a second read on your judgment, never a verdict over it. We are trying to follow what we teach, and we will not always get it perfectly right. Where you find us falling short of the standard we hold you to, we genuinely want to know.

At AI4LD, we believe in designing with integrity, teaching with clarity, and always putting the learner at the center.

Your Companion Tool for This Unit

This ebook has a companion: the **AI Literacy Curriculum Auditor**.

Where this ebook argues that AI literacy is four uneven dimensions rather than one prompting skill, the Auditor lets you test that argument against your own work. Paste in an existing or draft AI-literacy lesson, module, or course outline, and it scores the content against the four

dimensions this ebook develops: understanding, using, evaluating and creating, and ethical/human-centered engagement. Two flags are built in deliberately, because they are the two gaps the evidence base for this ebook keeps surfacing. The Auditor explicitly flags when the ethical/human-centered dimension is thin or absent, and when a module tests only text-based prompting while the stated learning objective also implies multimodal competence.

You do not need a finished module to use it. If all you have is a rough idea of what you might build (a topic list, a few sentences of intent), choose “Just an idea so far” and the Auditor gives you a design map rather than an audit: what your idea already carries, what to think about for each dimension, and where to start, drawn from the same evidence base as this ebook.

[AI Literacy Curriculum Auditor →](#)

Here is what the Auditor does and does not do. It checks your module against the four dimensions this ebook develops: which ones your content actually teaches, which are thin, and which are missing. It does not certify that a module is pedagogically sound, and it does not replace the judgment you bring to your own learner population, organizational context, and constraints. Treat its output the way you would treat a colleague’s structured second read: useful for catching what a first draft missed, not a verdict to defer to.

Introduction: AI Literacy Is Four Uneven Dimensions, Not One Prompting Skill

Why This Ebook Matters Now

You have almost certainly been asked to “build an AI literacy module”: for a client, for your own organization’s onboarding sequence, for a professional development track, or attempting to build your own AI literacy. And you have almost certainly reached, at least initially, for what is easiest to build and easiest to demonstrate, a set of prompting tips. How to write a clear instruction. How to iterate on a first draft. Maybe a short segment on hallucination.

This is a reasonable starting point. It is not, on the evidence, a complete one.

The core argument: AI literacy develops unevenly across four distinct dimensions: understanding how AI works, using AI effectively, evaluating and creating with AI output, and engaging with AI ethically and in human-centered ways. Left to develop on their own, through unstructured exposure and practice, these four dimensions do not develop at the same rate. The evidence converges on one dimension in particular being the one that spontaneous practice consistently fails to reach, the ethical and human-centered dimension. A module built entirely around prompting tips teaches, at most, one of the four dimensions well. It does not teach the other three, and it is least likely of all to teach the fourth.

This ebook treats prompting as one component of AI literacy, not as its substitute.

Much of the resistance to AI in learning and work is resistance to the unknown: not knowing what a system actually does, where it fails, or where human judgment still has to carry the weight. Literacy dissolves the unknown; reassurance does not. Understanding AI’s real capabilities and limits is Dimension 1 of this framework, not a separate argument from it. And taking that resistance seriously, treating multiple perspectives on this technology as worth engaging with rather than dismissing them, is itself the ethical, human-centered stance Dimension 4 asks of you.

A Critical Distinction: Prompting Skill Is Not AI Literacy

Prompting skill and AI literacy are not the same construct, and building a module that develops the first does not reliably produce the second.

Prompting skill is operational but also related to AI fluency, the capacity to produce a desired output from a generative AI tool through instruction, iteration, and refinement. Prompting strategies have also demonstrated variation over time as different models require different types of input to maximize their work and output. But the evidence base this ebook draws on treats AI literacy as a broader construct, organized around four dimensions: understanding what AI is and how it functions, using AI tools productively, evaluating and creating with AI-assisted output critically, and engaging with AI ethically and in ways that keep human judgment central (Sofkova Hashemi, 2026). Prompting skill sits almost entirely inside the second dimension. A curriculum that teaches prompting well and stops there has covered roughly a quarter of the construct it claims to be teaching.

The evidence for why this distinction is not purely academic is clear. In a workshop study of 28 postgraduate students collaboratively developing prompting skill over an extended, structured task, students' reflections showed strong development in understanding AI's mechanics and evaluating and creating with it, but the ethical and human-centered dimension remained comparatively underdeveloped, even in a well-designed learning environment with engaged, self-selected participants (Sofkova Hashemi, 2026). If the ethical dimension lags even under favorable conditions, it is very unlikely to develop on its own in a module that never addresses it directly.

What this framework does not claim. It does not claim that the four-dimension structure is the only valid way to organize AI literacy instruction, or that these five studies exhaust the evidence base for how each dimension develops. It does not claim that a module which formally addresses all four dimensions guarantees ethical or effective AI use by every learner who completes it. Curriculum design shapes the conditions for learning; it does not determine the outcome for every individual. What it does claim, on the evidence available, is narrower and more useful. Modules organized only around prompting tips are systematically incomplete, and the dimension most consistently missing is the one instructional designers are least likely to have been asked to build.

A Personal Connection

Everywhere I look in the L&D field today (conferences, webinars, social media), the emphasis still falls on the tools themselves rather than on a human-centered design perspective that explicitly considers AI literacy, ethics, transparency, and explainability. The field does seem to be coming around to these questions, finally, especially as large language models and other AI-embedded tools keep expanding what they can do. One of the most pressing needs of our time is holding the balance between using these tools and critically evaluating them, understanding their affordances, not just their outputs.

I have been on my own journey for several years now, trying to understand these complex computational environments and what our field can actually learn from them to improve learning and performance. A computer science colleague of mine at George Mason University has worked to get her students to sit with the ethical questions AI raises in the classroom. What is an AI-detector actually doing with the text you just submitted to it? Where does that text go once it's processed? Does it end up training the models being used to police it? The instructional design task is to increase awareness and identify the teaching moments that sharpen student's understanding of why the system might have produced a particular output and what that means for our instructional or training module design. After designing and delivering her own course on AI literacy, she put it this way: "The deliberate design of what AI literacy means inside a course, how a student forms a position, tests it, and learns to read a machine's output for what it gets wrong, is the actual work in front of us." Her caution for the field is worth sitting with too: "Things are moving very fast, and we have very little time to get AI literacy right" (Shehu, 2026).

What the research reviewed in this ebook offers is not a reason to abandon the tool-first starting point most of us default to, but a structure for what has to sit alongside it.

What You'll Learn

In the chapters ahead, we'll cover:

The AI Literacy Evidence Base: What recent educational research actually finds about how AI literacy develops: where it develops well on its own, where it doesn't, and the metacognitive risks and demands that curriculum design has to address.

The Four-Dimension Framework: A practical structure (understanding AI, using AI, evaluating and creating with AI, and ethical/human-centered AI) mapped against what the evidence shows about each, with the specific gaps named directly: the multimodal prompting gap, the functional-versus-structural understanding gap, and the field-wide absence of research on teaching the ethical dimension.

Designing an AI Literacy Module: Concrete design strategies, an audit checklist, and a worked scenario for building, or rebuilding, an AI literacy module that addresses all four dimensions deliberately, rather than assuming the three dimensions beyond prompting will develop on their own.

By the end of this ebook, you will have a framework for auditing any existing AI literacy module against the evidence, and a set of design moves you can apply to your own module this week.

AI4LD Learning Objectives

By the end of this AI4LD learning experience, you will be able to:

ANALYZE

Explain the four dimensions of AI literacy: understanding, using, evaluating/creating, and ethical/human-centered engagement. Distinguish prompting skill (a component of the “using” dimension) from the broader construct. Identify the specific gaps the evidence names in each dimension: the multimodal prompting gap, the functional-versus-structural understanding gap, and the underdeveloped ethical dimension.

EVALUATE

Critically evaluate an existing or draft AI literacy module against all four dimensions, not just prompting competence. Assess whether a module’s design assumes verification behavior and ethical reasoning will emerge spontaneously, or builds them in deliberately. Judge whether a module is calibrated to learner skill level, given the evidence that overreliance and under-utilization are both real risks at different points on that spectrum.

CREATE

Design or redesign an AI literacy module that deliberately addresses all four dimensions, with explicit attention to the dimension the evidence shows is most consistently underdeveloped. Apply the four-dimension audit to an existing learning environment to surface curriculum design decisions that were made by default. Develop a set of design principles for your own organizational context that reflects the framework introduced in this guide.

Chapter 1: The AI Literacy Evidence Base — What the Research Actually Shows

Why Start Here

Instructional designers building an AI literacy module rarely have the time to read five separate empirical studies before drafting a course outline. This chapter synthesizes these study results to inform your ID/LXD practice. It reviews four findings from recent educational research: the four-dimension structure of AI literacy that includes prompting as an evolving practice rather than a one-time skill, the gap between using AI and understanding AI, and the metacognitive question instructional and curriculum design needs to consider. Other current studies are referred to briefly under this structure acknowledging that AI literacy frameworks, even when aligned with different target audience groups are still in their infancy (Chee, Ahn, & Lee, 2025). With these considerations, chapter 2 then turns these exploratory findings and insights into instructional design recommendations for consideration.

AI Literacy Is Not a Single, One-Shot Skill — It’s Potentially Four Rich Aspects and Much More

The four-dimension structure this ebook uses (understanding AI, using AI, evaluating and creating with AI, and ethical/human-centered engagement with AI) comes from a competency development framework proposed in a review of thirty peer-reviewed studies (Ng, Leung, Chu, & Qiao, 2021) and incorporated into one educational research study this ebook also draws on (Sofkova Hashemi, 2026). The OECD frames AI literacy the same way, as a complex construct, not a simple practical skill or a body of knowledge to memorize (OECD & European Commission, 2026).

What that complexity looks like in practice is that the four dimensions don’t develop at the same rate, even when a learner is actively engaged and practicing all four. The clearest evidence for this comes from an exploratory workshop study of 28 postgraduate students’ self-documented prompting histories. Reflections aligned closely with three of the four dimensions, showing a growing understanding of how GenAI interprets prompts, increasingly strategic prompt crafting, and increasingly critical evaluation of outputs for creativity, coherence, bias, and alignment with intent (Sofkova Hashemi, 2026). The fourth dimension told a different story. Students showed some awareness of bias and representational issues, but that awareness stayed inside their prompting and evaluation habits rather than becoming its own explicit ethical reasoning; it “remained less explicitly developed” than the other three (Sofkova Hashemi, 2026). This is

graduate-student evidence, not workplace evidence, but the pattern it documents, with three aspects or dimensions compounding through practice but the fourth stays implicit, is exactly what instructional designers in corporate settings need to focus on.

Design implication: don't assume all four dimensions develop together just because a module includes hands-on AI practice. Design specific, goal-oriented exercises for the "using" dimension, and build in explicit attention to functionality, explainability, privacy, transparency, and ethics, rather than leaving the fourth dimension to develop as a byproduct of prompting practice. Planning deliberate reflection on the ethical/human-centered dimension across modules and activities is what closes the gap this evidence identifies.

Prompting as Epistemic Work, Not a Technical Instruction

A second finding reframes what "teaching prompting" should even mean, with more focus on the epistemic work, how we come to know and understand things, that happens in dialogue with large language models. Rather than treating prompting as a fixed technical skill you demonstrate once, the workshop study found that students' prompting strategies evolved over the course of the task: from simple one-shot requests toward iterative negotiation with the system, including strategies like asking the system to first articulate "what makes a good story" before building on that answer, and role-play framing to shape tone and perspective (Sofkova Hashemi, 2026). The author frames this evolution as **epistemic work**: a practice through which learners interpret, negotiate, and critically evaluate AI-generated content, not declarative content added on as a module they simply memorize.

Students in the same study also discovered something about the system's behavior that prompting guidance would not have communicated: a **structural rigidity** in the output, where the model tended to preserve the underlying structure of an early writing draft (its "skeleton") even under substantial revision requests, changing surface details more readily than deep structure (Sofkova Hashemi, 2026). That is a conclusion about system behavior that students had to discover through iterative practice, not something a one-time prompting lesson could hand them in advance.

Design implication: a prompting component built as a single lesson, rather than as iterative practice provided across a module, will not fully address the "using AI" aspect it is meant to cover. Epistemic work requires repeated cycles of attempt, observe, reflect, and adjust, not a one-time demonstration of a single type of prompt. Several prompting techniques, including input-output framing, role-based directions, and step-by-step reasoning such as Chain of Thought (CoT) or asking for contextual explanations first, appear in this research as more advanced approaches. The deeper implication is that as students refined and negotiated their

prompting, their reasoning shifted toward higher-level thinking about multimodal AI affordances, human-AI agency, authorship ethics, and the system's strengths and limitations (Sofkova Hashemi, 2026). This is the kind of growth a targeted, reflective, and explicit epistemic-dialogue approach can promote learning across the more complex aspects of AI literacy.

Using AI Well Doesn't Always Mean Understanding AI

Chiu (2025) separates AI literacy from AI competency, arguing they operate at different levels of knowledge and skill. Chiu distinguishes the “skillful application and optimization” of AI competency from the more nuanced description of AI literacy as “...interacting with AI responsibly - questioning outputs, recognizing limitations, and making informed decisions as consumers, citizens or professionals” (Chiu, 2025, p. 3225). Another line of studies complicates the assumption that operational fluency produces conceptual understanding of AI. In a study of 67 pre-service teachers practicing intercultural conflict scenarios with an AI-powered roleplay tool, graduate students used significantly more technical and structural language to describe the system, and connected their reflections to a broader range of professional-practice domains, than undergraduate students did, even though both groups used the tool the same way (Böhmer et al., 2026). Graduate students' professional experience may partially explain the difference. More pointedly, across the entire group, using the tool effectively (technique/tool-oriented engagement) and understanding how or why it worked (structural engagement) were only weakly correlated. The authors frame it the same way Chiu (2025) does. **Functional engagement with these tools does not automatically produce structural understanding of them or how and why outputs were generated** (Böhmer et al., 2026).

This finding is not a result of one study or use of one tool. Research that measures AI literacy using self-report and demonstrated performance, keeps finding the two don't always align. Zhang et al. (2026) emphasized how weak that alignment actually is, providing evidence that self-reported confidence and demonstrated performance were not related, holding across all four aspects, including hands-on use and conceptual understanding of AI. Confidence doesn't fail uniformly either: some learners overestimate what they can do, some underestimate it, and a few haven't had enough exposure to judge themselves accurately yet. Telling which is which requires an actual performance check, not a self-rating. The field's own established self-assessment scales build this same distinction into their structure, treating “using” and “understanding” AI as separate facets, not two names for the same competency (Carolus et al., 2023). Feeling confident using a tool and actually understanding it are different things to measure, because they're different things to develop.

Workforce-specific competency research makes clear why this distinction matters in professional practice. Understanding how and why an AI model might have arrived at a particular output is what makes bias identifiable, decisions explainable, and recommendations trustworthy, and choosing the right AI tool for a task depends on knowing where that tool's competence is limited, such as a generative model's real limits on quantitative analysis (Chee et al., 2025). Chee et al. (2025) also name the career-related competency this produces, not just using a specialized AI tool, but analyzing and evaluating its output well enough to explain the result to a colleague or client. That's the throughline this section's evidence keeps pointing to. AI literacy doesn't develop from operational practice alone.

The same study reframes the AI tool's instructional role through Vygotsky's "More Knowledgeable Other." A generative AI system can partially fill that role, but because it functions as a scaffold, the study argues this raises an open question about who is responsible for which parts of the learning process, as the instructor's own role shifts toward what the authors call, somewhat facetiously, an "accompanying prompt engineer" (Böhmer et al., 2026). In a second analysis, Chiu and Rospigliosi (2026) expand the instructor's (or ID/LXD's) responsibility, analyzing six global AI literacy frameworks across K-12, higher education, and workforce development contexts to identify shared threads: 1) a human-centric orientation moving beyond just technical skills; 2) an ethical stance addressing bias, privacy, fairness, transparency, and responsible use; 3) a focus on critical thinking, societal awareness, human values, and evaluating AI outputs critically; 4) understanding, application, and creation while weighing both benefits and potential risks with AI; among others. These frameworks, and others like them, point toward how we might begin to cultivate AI literacy across instructional and training audiences, beyond mere use of these tools.

Design implication: operational practice, however extensive, is not a substitute for direct instruction on how and why generative AI systems actually function, a human-centered AI literacy approach tailored to the audience's specific needs. Aspects related to using and more deeply understanding AI have to be designed and integrated deliberately. The evidence shows they don't reliably transfer from one to the other, and that holds across both classroom research and workplace contexts. An instructional module shouldn't assume hands-on practice will produce deep understanding of these complex computational systems, and it probably shouldn't substitute a "how comfortable are you with AI?" self-report item for either one also. Cultivating AI literacy as the complex construct it is means attempting to design deliberately for the how and why AI can augment human judgment and consider its impact across different settings, not just for use of the tool itself.

Metacognitive Laziness or Unsupported Metacognitive Demands: A Question for the Field

The fourth finding raises the issue that seems to make these curriculum and instructional design considerations necessary rather than optional. A review of 38 studies on GenAI-supported computational thinking instruction found a recurring pattern the field has referred to as **metacognitive laziness**: reduced critical engagement and reflective effort when learners over-rely on AI output, reported in 44% of the reviewed studies (Ouaazki et al., 2026). The same review identified a genuinely mixed evidence base, with 66% of studies reporting improved learning outcomes, 16% reporting decreases, and 16% no change in learning. This suggests that “AI-enhanced” is not, on its own, a guarantee of better learning; instructional design choices may more significantly influence the overreliance on and learning outcomes with AI.

The study behind the term itself is worth examining directly. In a randomized experimental study, 117 university students completed an essay-writing task with one of four types of support: ChatGPT, a human expert, a writing-analytics tool, or no additional tool (Fan et al., 2025). The ChatGPT group showed the strongest short-term improvement in essay scores, but knowledge gain, transfer, and intrinsic motivation were not significantly different across the groups, and the groups differed in the frequency and sequencing of their self-regulated learning processes. The authors’ caution is stated carefully. AI technologies such as ChatGPT may improve short-term task performance, but there is a risk that learners may become overly reliant without full engagement in the learning task, triggering what they refer to as “metacognitive laziness” (Fan et al., 2025). The authors also acknowledge a number of limitations to the study and overgeneralizing these results. A related concern from the learning-to-learn literature reaches a similar conclusion. Reliable external availability of information may lead learners to retain knowledge of *where* to find something rather than the information itself, producing what those authors call an “illusion of competence” (Schorr et al., 2026).

More current work in human-computer interaction, however, raises the important point that what looks like laziness may be better understood as unsupported demand. Drawing on psychology, cognitive science, and recent GenAI user studies, researchers have argued that the core struggles people experience with generative AI (crafting prompts, evaluating outputs, deciding how much to rely on them, integrating these tools into workflows) are metacognitive in nature. These systems impose metacognitive demands on learners as well as require self-awareness of generative AI’s impact on the targeted workflow with flexibility and adaptability (Tankelevitch et al., 2024). Related to this reading, over-reliance may not simply reflect reduced learner effort; it may be a predictable response to a technology that raises the metacognitive workload of a task without supporting it. The same work suggests these demands can be addressed by

integrating metacognitive support strategies overtly and by designing for explainability and customizability (Tankelevitch et al., 2024). This begins to shift the instructional design question from warning learners about cognitive laziness toward how to best scaffold the monitoring, evaluation, and reliance decisions that working with these tools seems to require.

The risk also appears to run in both directions. The computational-thinking review found overreliance concentrated among beginners, while more advanced learners showed the opposite pattern, under-utilizing GenAI tools on complex tasks where the tools could have genuinely extended and supported their capability, a tension the authors describe as calibration failure at both ends of the skill spectrum (Ouaazki et al., 2026).

Design implication: a module cannot assume that including AI automatically improves learning, and a single design may not serve learners at different skill levels and in different contexts. Guarding against over-reliance for beginners and under-utilization for advanced learners remain two different design problems, and an instructional designer needs to know which learner population he or she is designing for. The metacognitive-demands perspective suggests a more constructive framing of both: rather than only warning learners about metacognitive laziness, design the metacognitive support into the learning experience itself: structured opportunities for planning and task decomposition, reflection on when and why to rely on an AI output, and explicit attention to calibrating confidence in one's own judgment versus the system's. At the very least, this literature raises the possibility that metacognitive engagement is something an instructional or curriculum design can deliberately support and practice.

Taken together, these four findings point toward the same conclusion from four different angles: AI literacy is multidimensional, uneven in how it develops, not reliably self-taught, and subject to metacognitive demands that deliberate design seems able to support. Chapter 2 turns this evidence into a structure you can try out to design with directly.

Chapter 2: The Four-Dimension Framework — A Structure for Instructional Designers and LXDs

Using the Framework as a Design Lens

The four aspects or dimensions below are not sequential steps. A learner does not “finish” understanding before moving to using AI. They develop in parallel, unevenly, and the evidence in this chapter names specifically where and why each one tends to fall short without deliberate instructional design. For each dimension, we cover what it is, what the evidence shows about how it develops, where spontaneous practice may consistently fall short, and the design implications that follow.

One more design principle cuts across all four dimensions before we get into them individually. For a workforce audience, AI literacy is field-specific, not generic. A systematic review of AI literacy competency research finds that once training moves into a professional setting, the competency that matters shifts from broad AI knowledge toward interpreting and applying AI tools within a specific career path, including error detection and AI-informed decision-making in that field's own terms (Chee et al., 2025). Practically, that means the examples and practice tasks in each dimension below should come from your learners' actual work, not generic exercises borrowed from an education setting.

Dimension 1: Understanding AI

WHAT IT IS

Understanding AI is conceptual and structural literacy. It means knowing what a generative AI system is, broadly how it produces output, and what its genuine capabilities and constraints are. Chiu (2025) draws this same line between AI competency, the “skillful application and optimization” of a tool, and AI literacy, “interacting with AI responsibly - questioning outputs, recognizing limitations, and making informed decisions as consumers, citizens or professionals” (p. 3225). This dimension sits on the literacy side of that distinction. It is not a requirement that learners understand model architecture at an engineering level. It is the literacy that lets a learner reason about *why* an AI tool behaves the way it does, rather than only observing *that* it behaves a certain way.

WHAT THE EVIDENCE SHOWS

Functional engagement with AI and structural understanding of it are not the same thing, a point Chapter 1 already made. Less obvious is how the gap survives even when a learner watches an AI explain itself. Palod, Biswas, and Kambhampati (2026) found that when an AI shows its reasoning or offers a post-hoc explanation, people trust the answer more, regardless of whether the answer is actually correct. The explanation is persuasive, not informative. The one condition that helped participants actually tell right answers from wrong ones was a dual explanation, arguments for and against the AI's position, presented together. Left with a single, confident-sounding explanation, learners read fluency as a signal of soundness, mistaking exposure to an AI's stated reasoning for genuine understanding of how it arrived there.

WHERE SPONTANEOUS PRACTICE FALLS SHORT

This is the **functional-versus-structural understanding gap**: using an AI tool effectively does not automatically produce an accurate mental model of how or why it works. A learner can become operationally skilled at prompting (through the kind of iterative, dialogic practice described in Chapter 1) without ever developing a structural understanding of what is actually happening when the system generates a response. The general pattern holds across very different methods: an exploratory qualitative study of teacher trainees (Böhmer et al., 2026), psychometric scale design (Carolus et al., 2023), and direct self-report-versus-performance comparisons (Zhang et al., 2026) all find the same split. Palod et al.'s (2026) finding shows why the gap is so hard to close through use alone. The one intervention that looks like it should teach understanding, watching an AI explain itself, instead teaches false confidence. A learner who has seen an AI “show its work” can walk away feeling like they understand the system better, with no actual gain in their ability to tell a sound answer from an unsound one. Left to develop through use alone, this dimension is the one most likely to be shallow even in learners who otherwise seem fluent, and explanation exposure by itself will not fix that. Naming this dimension isn't unique to this ebook's model, either: Chiu and Rospigliosi's (2026) comparison of six global AI literacy frameworks found a human-centric orientation moving beyond technical skills, and a focus on evaluating AI outputs critically, as important shared threads across independently developed frameworks.

DESIGN IMPLICATIONS

Teach the mechanism directly. Don't assume it will be inferred from practice. Build a module component or integrate into and across existing modules that explicitly address *how* generative AI produces output (in accessible, non-engineering terms), what a context window is and why it constrains what a system can respond to, and why outputs can be fluent and confident without being accurate. Don't rely on hands-on practice alone to produce this

understanding. Across a qualitative study of teacher trainees, a validated measurement scale, and direct self-report-to-performance comparisons, the pattern is the same and worth attention (Böhmer et al., 2026; Carolus et al., 2023; Zhang et al., 2026).

Don't let the AI's own explanation stand in for understanding. Palod et al. (2026) found that reasoning traces and post-hoc explanations increase trust in an AI's answer whether or not that answer is correct, and that only a dual explanation, arguments for and against the AI's position, actually helped people tell the difference. A module that shows learners an AI "explaining itself" and stops there is teaching them to trust fluent explanations, not to evaluate them. Build in a deliberate second step: ask learners to argue the case against the AI's stated reasoning before accepting it, or pair every AI explanation with a genuine counter-argument for comparison.

Practical example: A financial services onboarding program includes an AI literacy module built entirely around hands-on prompting exercises with the firm's internal assistant. After completing the module, new analysts can produce useful outputs quickly but describe the tool in almost entirely social, "it just knows" terms when asked how it works, with no structural vocabulary for its constraints. Adding a short, explicit "how this system actually works" segment before the hands-on practice (covering training data, pattern generation, and the fact that fluency is not the same as accuracy) measurably shifts how analysts describe the tool's limitations in a follow-up discussion. The program adds a second habit on top of that: whenever the assistant explains its own reasoning for a recommendation, analysts are prompted to draft the strongest counter-argument before accepting it, not just read the explanation and move on. Analysts who only got the mechanism segment could describe how the tool worked but still took its self-explanations at face value; analysts who also practiced arguing the other side caught materially more cases where a confident-sounding recommendation was actually wrong, without working any slower.

Dimension 2: Using AI and multimodality

WHAT IT IS

Using AI is operational literacy. It is the practical, evolving skill of interacting productively with generative AI tools across the modalities a learner's role actually requires. This includes prompting, but it is broader than prompting. It includes recognizing and adapting to system behavior, and applying strategy across different output types. The U.S. Department of Labor's national AI literacy framework names similar concerns as one of its five foundational content areas, prompting or directing AI effectively, distinct from understanding

how AI works and evaluating AI outputs (U.S. Department of Labor, 2026). This dimension isn't this ebook's own division of the field; it matches federal workforce guidance built for the same adult, professional audience.

WHAT THE EVIDENCE SHOWS

Prompting is iterative, negotiated skill-building, not a one-time demonstration, as Chapter 1 already showed (Sofkova Hashemi, 2026). That same evidence also points to a limit on the iteration. Skill built through text prompting doesn't automatically carry into other modalities, even for learners who have genuinely developed the "using AI" skill in text.

WHERE SPONTANEOUS PRACTICE FALLS SHORT

This is the **multimodal prompting gap**: text-prompting competence does not necessarily transfer to image-prompting competence. Students who had become skilled at text prompts found image generation markedly harder to prompt effectively. Narrative and emotional context that was implicit in their text prompts did not automatically carry over to their image prompts and had to be explicitly re-specified (Sofkova Hashemi, 2026). Most modules that teach "prompting" teach only text prompting, which means they teach, at best, half of the operational dimension for any role that also requires visual, audio, or other non-text output. This isn't isolated to one classroom study, either. Workplace AI-training guidance reaches the same conclusion from a completely different direction. A professional guide to multimodal AI use names context bridging, carrying meaning across formats without losing it, as one of three core patterns any professional needs, alongside cross-modal understanding and cross-modal generation, and frames effective multimodal AI use as inherently iterative, not a one-shot prompt (Simms, 2025). What the study found students struggling to do without instruction is exactly what workplace AI training is now teaching professionals to do deliberately, and it is a small piece of a much larger operational skill set, not the whole of it, which is the point this ebook is making by moving past prompt engineering in the first place.

DESIGN IMPLICATIONS

Sequence practice across at least two modalities if the target role requires more than text output. If learners will ever need to prompt for images, audio, or other non-text formats in their actual work, a text-only prompting module has not covered the "using" dimension completely. It has covered a subset of it that the evidence shows does not transfer.

Teach re-contextualization as an explicit, named skill. Because context that is implicit in a text prompt does not automatically carry into an image prompt, a module should teach learners to actively check what context needs to be restated, the skill workplace AI-training guidance

names as context bridging, rather than assuming the system will infer it (Sofkova Hashemi, 2026; Simms, 2025).

Practical example: A marketing team's AI literacy module trains staff exclusively on text-based copy generation. Six weeks later, the same staff are asked to generate accompanying visuals for a campaign and produce inconsistent, generic images that don't match the tone established in their text drafts. The gap is not a tooling problem. It's an untaught skill. Adding a single multimodal practice sequence, where learners take a text prompt they've already refined and explicitly re-specify its tone and context for an image prompt, closes most of the gap in a subsequent trial.

Worked example: re-specifying the prompt. Here's what "explicitly re-specify tone and context" actually looks like in practice, using the same campaign.

Text prompt (already refined for the copy): "Write a two-sentence social ad headline and subhead for a budgeting app, aimed at first-time freelancers who feel anxious about inconsistent income. Reassuring but not condescending. No corporate finance jargon."

Naive image prompt (the tone and audience context left implicit): "Create an image for a budgeting app ad."

Re-specified image prompt (the same context carried forward explicitly): "Create an image for a budgeting app ad aimed at first-time freelancers who feel anxious about inconsistent income. A single person in a calm, everyday setting: kitchen table, laptop, coffee, warm natural lighting. Approachable and reassuring, not corporate or slick. No stock-photo stiffness."

The first image prompt produces something generic because none of the audience or tone information survived the switch. The second one names it directly. Carrying that context across the modality switch without losing it is what Simms (2025) calls **context bridging**, one of three patterns his workplace guide names for professional multimodal AI use, alongside cross-modal understanding and cross-modal generation.

A last step puts the second of those patterns to work. The learner asks the AI to check the finished image against the original ad copy, whether the visual actually reads as calm and reassuring against the same tone the text sets. That's **cross-modal understanding**, connecting information across formats instead of treating each one in isolation, rather than assuming a prompt that reads well in one modality has automatically succeeded in another.

That's the whole skill this dimension asks a learner to practice, not a new technique, just the discipline of checking what a new modality needs restated, and then checking that the two formats actually agree.

Dimension 3: Evaluating and Creating with AI

WHAT IT IS

Evaluating and creating with AI is the critical-application dimension. It means judging the quality, accuracy, and appropriateness of AI output, and integrating that output into original analytical or creative work rather than accepting it at face value.

WHAT THE EVIDENCE SHOWS

Where this dimension develops well, it looks like specific, nameable behaviors. In the exploratory workshop study, students shifted from basic interaction toward actively refining their prompts: improving alignment with their own intentions, interrogating outputs, requesting explanations, and adjusting inputs to compensate for the system's limitations. One concrete instance was refining prompts when image outputs defaulted to stereotypical representations. The study reads that dissatisfaction-driven refining loop as growing capacity to critically evaluate and create with AI (Sofkova Hashemi, 2026). Prompting emerged there as a strategic practice, where effectiveness depended on balancing specificity and openness, not simply adding detail. And it looks like **trust-but-verify**: treating AI output as a starting point to check, not a finished answer to accept. In a case study of undergraduates, students who reported finding AI genuinely useful consistently paired that account with active verification practices, corroborating AI output against course materials, requesting citations, cross-checking sources. This wasn't incidental. Network analysis of their interview discourse linked the verification behavior systematically to learning-support discourse, functioning as a form of "competence protection" (Isaeva et al., 2026). The same study found students situated AI as one node in a broader learning network rather than a substitute for their instructors, explicitly naming their human instructors as sources of motivation, ethical guidance, and emotional support that AI could not replicate (Isaeva et al., 2026).

WHERE SPONTANEOUS PRACTICE FALLS SHORT

Trust-but-verify is a real, teachable behavior. But it is not guaranteed. Chapter 1 already covered the underlying risk, metacognitive laziness, the reduced critical engagement that shows up when learners lean on AI output rather than evaluate it. The evidence adds something further here. The risk isn't symmetric. Beginners are the ones most likely to skip verification and over-rely; more advanced learners often swing the opposite way, under-using AI on complex tasks it could genuinely help with. The authors frame this as a **"Zone of No Development"**: too much unstructured support collapses a novice's need to build independent judgment, while too little contextualized support leaves an advanced learner's capability

underused (Ouaazki et al., 2026). Either failure looks the same from the outside, disengaged use, but the fix runs in opposite directions, which is exactly why a single, one-size-fits-all approach to AI access doesn't work.

DESIGN IMPLICATIONS

Teach verification strategies explicitly, as curriculum content. Don't assume trust-but-verify emerges on its own. The specific behaviors identified in the evidence (requesting citations, cross-checking against independent sources, corroborating against course or organizational materials) can be taught directly as a skill, rather than left to develop as an incidental byproduct of AI use (Isaeva et al., 2026).

Calibrate the level of AI access to the learner's skill level. For beginners, the evidence supports explaining risks directly and favoring tutor-style, non-answer-giving prompting patterns that require the learner to attempt work before receiving AI assistance. For advanced learners, the evidence supports more open-ended, context-rich integration (giving the tool fuller access to the task context rather than constraining it) because under-utilization, not overreliance, may negatively impact learning for this target audience (Ouaazki et al., 2026).

Practical example: A software company gives every support engineer, junior and senior alike, the same open-ended AI assistant for diagnosing customer-reported bugs. Junior engineers accept the assistant's first plausible diagnosis and close tickets that reopen a week later, unverified against the actual system logs; senior engineers, given the same tool with the same instructions, barely use it, because its generic suggestions don't match how they already triage a real incident. Redesigning the workflow with two tracks, a structured, verification-required diagnostic checklist for juniors and full, context-rich access for seniors to use as they see fit, cuts reopened tickets on the junior track and gets genuine adoption on the senior one, without changing the underlying tool.

Dimension 4: Ethical and Human-Centered AI

WHAT IT IS

The ethical and human-centered dimension covers bias awareness, integrity and transparency in AI use, appropriate human-AI role boundaries, and consideration of AI's effects on other people. This is the dimension most directly connected to whether AI use in a learning or work context remains accountable to human judgment rather than displacing it.

WHAT THE EVIDENCE SHOWS

This is the dimension these research findings suggest does not develop without instructional support. Even the exploratory postgraduate workshop study's fully engaged students showed the pattern Chapter 1 already named, a dimension that "remained less explicitly developed" than the other three (Sofkova Hashemi, 2026). Their encounters with ethical material, stereotyped defaults in AI-generated images, inconsistent copyright refusals across tools, shaped their prompting and evaluation, but never became explicit ethical reasoning about what should be generated or delivered.

The strongest version of this finding comes from outside the prompting literature entirely. A review synthesizing 21 studies on learning-to-learn competencies in the GenAI era built a three-layer framework that explicitly includes ethical dispositions as one of five core dimensions, and found it the least developed in the literature the review covers. More importantly, the authors found limited evidence of studies integrating learning-to-learn principles with generative AI, especially addressing the ethical dimensions (Schorr et al., 2026).

Another study supports this finding in its review and generation of design guidelines to include designing with digital ethics and integrity (addressing bias, privacy, transparency) (Ouaazki et al., 2026). These researchers found that these guidelines were not reliably addressed by existing GenAI education efforts.

A design-based study offers a glimpse into what embedding Human-Centered AI (HCAI) principles into instructional design contexts might encompass. The study investigated how seven teacher educators evaluated AI outputs, dealt with bias, and established policies for their own contexts across three cycles of needs analysis, prototyping, and refinement. The researchers concluded that AI literacy and HCAI meet in the enacted instructional design process within unique contextual and sociotechnical constraints, and should not be treated merely as skills-based training (Baran et al., 2026). These authors promote developing the professional judgment "...to interrogate AI outputs, justify use decisions, document limitations, and reflect on consequences" that at a broad level may also reflect similar challenges that L&D professionals face with different contextual factors, policies and constraints (p. 15). At the very least, this exploratory design research study raises the important point that in human-centered design, AI transparency, explainability and human-in-the-loop principles can be facilitated and enacted through human collaboration, modeling and sociotechnical interaction with AI.

WHERE SPONTANEOUS PRACTICE FALLS SHORT

The evidence in this area for the L&D audience is thin. That's not because the field lacks ethical frameworks or guidance. AI ethics scholarship, institutional principles, and dedicated research centers already exist in abundance, from explainability and FATE frameworks for educational AI

(Khosravi et al., 2022) to human-centered AI as a design discipline (Xu et al., 2023), with responsible-AI practices in educational research now systematically reviewed in their own right (Fu & Weng, 2024). Nor is the ethical dimension not being taught or researched. A recent systematic review of K-12 AI ethics education maps sixty-eight studies with pedagogical designs and assessed learning outcomes across cognitive, affective, and behavioral domains (Ma et al., 2025). What's missing is uncovering similar findings in adult and professional AI-supported learning contexts for andragogical (adult learning focused) instructional designs. In the empirical GenAI literature reviewed here, there is little evidence to follow for instructional models that turn ethical guidance into taught, practiced, assessed reasoning. The research exploratory workshop evidence shows what happens without that design. As the Sofkova Hashemi (2026) study suggested, encounters with bias and representational issues do real work, driving prompt refinement and critical evaluation, but that growth registers in the evaluating-and-creating activities while the ethical and human-centered dimension seems to stay less developed. This single exploratory study cannot settle the question, but it illustrates the nature of the gap.

DESIGN IMPLICATIONS

Make this dimension explicit, not separated from the sociocultural practice of the context you're designing for. A design-based study working directly with practitioners on this exact question found that ethical modeling has to be enacted within that practice, under real contextual and sociotechnical constraints, rather than delivered as separable skills-based training (Baran et al., 2026). No validated pedagogy yet exists for this dimension in adult and professional contexts, so the safest instructional design approach is a conservative one: structured reflection, scenario-based discussion, and explicit naming of bias and integrity issues when they occur, built into the same activities and workflows that teach the other three dimensions rather than isolated into a separate module. Draw on the ethical frameworks that do exist, FATE and explainability-centered design (Khosravi et al., 2022), human-centered AI (Xu et al., 2023), and the pedagogical designs reviewed in K-12 AI ethics education (Ma et al., 2025), rather than assuming any single instructional technique is proven effective in this context (Schorr et al., 2026).

Build structured reflection around the encounters your module already creates. If your module includes hands-on AI practice, learners will very likely encounter something ethically relevant: a biased default, an inconsistent policy, a confidently wrong output presented as fact. The exploratory workshop evidence suggests engaged learners will notice these moments and treat them as prompting problems to solve, refining until the output looks right, which is genuine

evaluative skill (Sofkova Hashemi, 2026). The design task is to sharpen that noticing. A deliberate reflection prompt turns “I fixed the biased image” into “why did the system default there, and what does that mean for what we generate?”

Name the human role explicitly, rather than leaving it implicit. Learners in one study located ethical guidance, motivation, and emotional support specifically with their human instructors, not with AI (Isaeva et al., 2026). That doesn't mean AI has no role in surfacing ethical questions or flagging risk; a tool can genuinely support that work. But the guidance itself should stay visibly anchored to a human instructor or mentor who can perceive the nuances and implications, not be positioned as something the AI tool delivers on its own.

Practical example: A financial-services company's L&D team pilots an AI drafting tool for compliance-training content. Instead of treating tool operation as the whole module, the design team builds judgment into the workflow itself: reviewers must interrogate each AI draft against the regulation it's meant to cover, justify why it was accepted, edited, or rejected, document any limitation they noticed, and flag what could go wrong downstream if that limitation goes unaddressed. New trainees review a set of these documented decisions before reviewing AI drafts on their own. The ethical dimension isn't a separate unit bolted onto tool training. It's demonstrated in how the team's own review process works, under the company's actual compliance constraints.

A Summary: The Four-Dimension Map

Dimension	What It Covers	Where the Evidence Says It Falls Short	Key Design Move
Understanding AI	How generative AI works, conceptually and structurally	Functional-vs-structural gap: using AI well doesn't produce understanding of it (Böhmer et al., 2026)	Teach the mechanism directly; don't rely on inference from practice
Using AI	Operational, iterative prompting skill	Multimodal gap: text-prompting skill doesn't transfer to image prompting (Sofkova Hashemi, 2026)	Sequence practice across modalities the role actually requires
Evaluating & Creating with AI	Critical judgment of AI output; trust-but-verify behavior	Metacognitive laziness for beginners; under-utilization for advanced learners (Ouaazki et al., 2026)	Teach verification explicitly; calibrate access to skill level
Ethical & Human-Centered AI	Bias awareness, integrity, human-AI role boundaries	"Less explicitly developed" even among engaged students (exploratory study: Sofkova Hashemi, 2026); zero studies found teaching it directly (Schorr et al., 2026)	Make explicit within practice, not as a separate module; build structured reflection around encounters

Chapter 3: Designing an AI Literacy Module

The Core Design Question

Across all four dimensions, one design question determines whether a module is teaching AI literacy or teaching only prompting:

Does this module deliberately address all four dimensions, or does it teach prompting and assume the other three will follow?

The evidence in Chapters 1 and 2 gives a consistent answer to what happens when a module assumes rather than designs. The ethical dimension in particular does not reliably develop on its own, and the multimodal and structural-understanding gaps persist even in well-designed, engaged learning environments. This chapter offers four design strategies for closing those gaps deliberately.

Designing deliberately also means designing with the people the module is for, not in isolation from them. The clearest evidence for this comes from the ethical dimension. Baran et al.'s (2026) finding is that ethical modeling has to be enacted within the sociocultural practice of the people you're designing for, not delivered as separable training. But the design-research method behind that study may promote more than the finding itself. It is also not a new principle in design research generally. My own Integrative Learning Design Framework has argued for over two decades that design and research are mutually constitutive, with practitioner involvement built into every phase rather than added at the end (Bannan-Ritland, 2003; Bannan, 2013; Bannan, Kelly, & Baek, 2008). Baran et al.'s finding is a recent, AI-specific instance of a much older design-research principle and commitment. Each strategy below works best when it is checked against, and where possible built with, the actual people who will use the module considering what they already do, what constraints genuinely shape their work, and what they would need to see modeled in a participatory design stance.

Design Strategy 1: Build Multimodal Practice Into Every Module Where the Role Requires It

The most common single omission in prompting-focused modules is testing only text prompting when the target role also requires other modalities. The evidence is specific about why this matters. Text-prompting skill does not transfer to image-prompting skill, because context implicit

in a text prompt is not automatically present once the modality changes (Sofkova Hashemi, 2026).

Practical design approach: Identify every output modality the learner’s actual role will require (text, image, possibly audio, or code) before designing the prompting component, by asking people who currently do the role rather than assuming from a job title. If more than one modality is required, build at least one exercise that takes a learner from a refined text prompt to an equivalent prompt in the second modality, and make the re-specification step explicit rather than assumed.

Critical consideration: Not every role requires multimodal competence. A module for a role that genuinely only requires text-based AI interaction does not need an image-prompting component added for its own sake. The design move is to match modality coverage to the role’s actual requirements, not to maximize modality coverage universally.

Design Strategy 2: Teach the Mechanism Alongside the Content or Skill

Hands-on prompting practice, however extensive, does not reliably produce structural understanding of how generative AI works. The two were only weakly correlated even among learners with substantial exposure (Böhmer et al., 2026).

Practical design approach: Add a short, explicit “how this actually works” component before or alongside hands-on practice: covering, in accessible terms, how outputs are generated from patterns in training data, why fluency and confidence are not evidence of accuracy, and what a context window is and why it constrains what the system can respond to. This does not need to be technically deep to close the gap the evidence identifies; it needs to exist at all, explicitly, rather than being assumed to emerge from use.

Example: A retail company’s AI literacy module for store managers originally consisted entirely of a single 20-minute hands-on session with the company’s AI scheduling assistant. Managers became comfortable using the tool quickly but, in a follow-up survey, several described it in language indistinguishable from a human scheduler (“it knows who’s available”). Adding a five-minute explicit segment on how the tool generates scheduling suggestions from historical data patterns, before the hands-on session, shifted how managers described the tool’s limitations without adding meaningful time to the module.

Design Strategy 3: Make Verification a Taught, Calibrated Skill

Trust-but-verify is real and teachable, but it is a specific set of behaviors (requesting citations, cross-checking against independent sources, corroborating against known-good materials), not a general disposition that develops on its own (Isaeva et al., 2026). At the same time, the risk profile differs by skill level. Beginners are more prone to overreliance and metacognitive laziness, while advanced learners are more prone to under-utilizing AI on complex tasks where it could genuinely help (Ouaazki et al., 2026).

Practical design approach: Teach the specific verification behaviors directly, as a checklist or worked example, not an abstract instruction to “think critically.” Then calibrate access by skill level, informed by how verification actually happens in the target workflow (which sources are trustworthy in that domain, what corroboration is realistic given actual time constraints) rather than a generic checklist: for beginners, favor tutor-style prompting patterns that require an attempt before AI assistance is available; for advanced learners, favor more open-ended access with full task context, since the evidence suggests constraining experienced learners’ access is more likely to produce under-utilization than to protect against overreliance.

Critical consideration: This is a differentiated target-audience design decision, not a single approach. A module built for a mixed-experience audience could either branch by skill level or make an explicit, instructor-facing decision about what is optimal for a specific audience.

Design Strategy 4: Make the Ethical Dimension Explicit Within Practice, Not Apart From It

This is the strategy the evidence most consistently supports as necessary (making this dimension explicit rather than assumed), though the design-based evidence below changes what that form should look like. The ethical aspects of AI literacy are not prevalent in the literature, particularly related to workforce contexts and does not fully address how to actively teach the important ethical dimensions of AI literacy (Schorr et al., 2026); in the exploratory workshop evidence, encounters with bias fed prompt refinement and critical evaluation rather than explicitly developed ethical reasoning (Sofkova Hashemi, 2026); and a parallel scoping review names digital ethics and integrity as a design guideline specifically because it is not otherwise reliably addressed (Ouaazki et al., 2026). Working directly with practitioners designing their own professional learning, a design-based study found the opposite of a stand-alone unit. Ethical modeling has to be enacted within the sociocultural practice of the actual instructional-design workflow, not delivered as separable skills-based training (Baran et al., 2026).

Practical design approach: Make the ethical dimension an explicit, designed-for objective, not a hoped-for byproduct of hands-on practice, and attempt to design it within the sociocultural practice of the context you're building for, not apart from it in its own module. Build structured reflection prompts around the ethical encounters your other exercises will likely already create (biased defaults, inconsistent policy behavior, confidently wrong output), and, where the workflow allows it, model the judgment itself. Let learners see how an experienced practitioner interrogates an AI output, justifies a use decision, or documents a limitation, not just discuss the concept in the abstract (Baran et al., 2026). Because no single validated pedagogy exists yet for this dimension specifically despite increasing interest by the researchers in the field, attempt to favor critical interrogation and discussion, scenario-based learning contexts, and close examination of AI output and reasoning.

Example: A hospital system has its own clinical documentation team help design the AI literacy module for its AI drafting tool, rather than building the module first and testing it on them afterward. In a short needs-analysis session, the clinicians surface constraints the design team hadn't considered: some patient details that can never be entered into the tool under HIPAA, which AI-drafted language needs mandatory physician review before it reaches a chart, and where a confident AI tone can mask a data gap a clinician would normally flag. The module is built directly around those constraints. The ethical judgment work is part of the workflow itself, not added on after a demo. A second short cycle with the same clinicians catches a HIPAA edge case the first design missed before the module is used.

A Fifth Consideration: Designing for a Skeptical Audience

The four design strategies above assume a learner is willing to engage with the module. That assumption increasingly may not be safe to make. Commentary on AI in education, higher education's own response to it, and AI's effect on workers describes a real and reasonable skepticism by learners: concern about job security, about cognitive offloading, about whose interests a given AI deployment actually serves. This commentary — Rebecca Winthrop, Seth Harris, and George Siemens, each interviewed by Allison Dulin Salisbury for *The Humanist* (Winthrop, in Salisbury, 2026; Siemens, in Salisbury, 2025; Harris, in Salisbury, 2025) — sits outside the peer-reviewed literature this ebook otherwise draws on, and is named here as field perspective rather than evidence. Siemens in particular argues the deeper issue is institutions adopting AI without deciding what knowledge and skills actually stay stable. This is a version of “why this, why now” aimed at the institution, not just the individual learner, that the design approach below borrows directly.

Practical design approach: Don't design past this skepticism. Design with it in mind. An explicit "why this, why now" framing early in a module, one that names the legitimate version of these concerns rather than assuming they don't exist, does more to build trust than a module that proceeds as if adoption were self-evidently good. Where the context allows it, giving learners some real input into how an AI tool gets used in their own practice, not just being told it will be used, treats this warranted skepticism as a starting position to design with, not resistance to route around.

Critical consideration: This is not a claim that skepticism about AI is always well-founded, or that it should always be accommodated. Some learner resistance reflects a few gaps this ebook's four dimensions attempts to address (misunderstanding how the tool works, not having seen verification behavior modeled, etc.). The design task is detecting these and attempting to distinguish and honor them, and not treating all resistance as either fully valid or fully solvable by better curriculum or instructional design.

Apply the Four-Dimension Audit to Your Module

Before an AI literacy module goes live, or as a structured review of an existing one, apply this audit. For each dimension, check whether the module addresses it deliberately or leaves it to emerge on its own.

Audit Question	What to Check
Does the module explain how AI actually produces output, not just how to prompt it?	Is there a distinct “how this works” component, or is understanding assumed to follow from practice?
Does the module cover every modality the learner’s role actually requires?	If the role requires image, audio, or other non-text output, is that modality practiced explicitly, not assumed to transfer from text?
Does the module teach specific verification behaviors, or just tell learners to “think critically”?	Are citation-checking, cross-referencing, and corroboration taught as concrete, practiced skills?
Is AI access calibrated to learner skill level?	Does the module risk overreliance for beginners, under-utilization for advanced learners, or both, by using one design for a mixed-skill audience?
Is the ethical/human-centered dimension made explicit within the module’s own activities, designed into practice rather than left implicit or a separate module?	Or does it assume ethical reasoning will emerge incidentally from hands-on practice?

Running this audit, or using the AI Literacy Curriculum Auditor described at the front of this ebook, can surface which of the four dimensions a module was designed around and which were assumed. The audit does not require deep AI expertise. It requires the same instructional design judgment you bring to any curriculum or instructional design review.

An LXD Scenario: The AI Literacy Module That Looked Complete

Consider an AI literacy module built for a professional services firm’s new-hire onboarding. The module includes a well-produced two-hour session: an overview of the firm’s approved AI tools, a hands-on prompting workshop with realistic client-facing scenarios, and a short closing quiz on prompting best practices. Post-session survey scores are strong. New hires report feeling confident using the tools.

Three months later, a client-facing team uses the firm’s AI assistant to draft a set of marketing visuals for a client pitch. The images default to a narrow, homogeneous depiction of the client’s stated target audience. It’s a default nobody on the team thinks to question, because nothing in their onboarding ever asked them to notice or interrogate a default like this. The pitch goes out;

the client raises the issue directly. The firm's after-action review finds the gap immediately. The onboarding module taught prompting well, and taught nothing about how to evaluate AI-generated defaults for the biases they may encode.

What went wrong, by dimension: the module covered “using AI” thoroughly. The hands-on workshop was genuinely well designed for prompting practice. It touched “understanding AI” lightly, through a brief tool overview, without addressing how or why outputs are generated the way they are. It did not cover multimodal prompting at all, despite the client-facing team's actual work requiring visual output regularly. And it had no ethical/human-centered component whatsoever. Nothing was designed to turn a new hire's encounter with a biased default into a named, questioned decision, rather than a prompting inconvenience to fix or, worse, miss entirely before it reached a client.

The redesigned module keeps the same hands-on prompting workshop, because it was genuinely effective for what it covered. It adds a short mechanism segment before the workshop begins. It adds a multimodal practice sequence, since visual output is a real part of the client-facing role. And it adds a structured, ten-minute ethical debrief immediately after the hands-on session, built directly around whatever biased or unexpected defaults the workshop's own exercises surfaced that day, turning a live example into the discussion, rather than a hypothetical one.

The same tool access and the same hands-on workshop now address four dimensions instead of one.

Reflection Prompts for Your Practice

As you work through the ideas in this guide, consider the following questions in relation to your current or planned AI literacy modules.

On understanding AI: Does your current module explain how AI generates output, or does it move directly into hands-on practice? If a learner finished your module and described the tool as something that “just knows” things, would you notice? Could this be an opportunity for instructional design?

On using AI: What modalities does your learners’ actual role require? Does your module’s prompting practice cover those modalities, or only text, and if only text, is that a deliberate scope decision or an unexamined default? How can you increase literacy about the consideration and differentiation of multiple modalities and AI?

On evaluating and creating with AI: Does your module teach specific verification behaviors (requesting citations, cross-checking sources), or does it assume learners will develop these on their own? Is your module designed for a single skill level, or does it need to calibrate access differently for beginners and advanced learners?

On the ethical/human-centered dimension: Does your module have a distinct component addressing bias, integrity, and human-AI role boundaries, or does it assume ethical reasoning will emerge from hands-on practice? If your hands-on exercises are likely to surface a biased default or a confidently wrong output, does your module have a designed moment to discuss it?

On the module as a whole: If you ran the four-dimension audit against your current module today, which dimension would score weakest? Is that the dimension the evidence would predict?

Consider these prompts as the beginning of a curriculum review conversation with your team, or as the basis for a session with the AI Literacy Curriculum Auditor described at the front of this ebook. The goal is not to identify failures but to surface design decisions that may have been made by default rather than by choice, and to reclaim the ethical and human-centered dimension in particular as a design decision rather than an assumption.

References

- Bannan, B. (2013). The Integrative Learning Design Framework: An illustrated example from the domain of instructional technology. In T. Plomp & N. Nieveen (Eds.), *Educational design research Part A: An introduction*. SLO (Netherlands Institute for Curriculum Development).
- Bannan, B., Kelly, A. E., & Baek, J. Y. (2008). The “compleat” design experiment: From soup to nuts. In A. E. Kelly, R. A. Lesh, & J. Y. Baek (Eds.), *Handbook of design research methods in education: Innovations in science, technology, engineering, and mathematics learning and teaching*. Routledge.
- Bannan-Ritland, B. (2003). The role of design in research: The Integrative Learning Design Framework. *Educational Researcher*, 32(1), 21–24.
- Baran, E., Dilek, M., Ziba, M., & Xiao, X. (2026). Human-centered AI for teacher educators: Designing professional learning for critical AI literacy. *Computers and Education Open*. Advance online publication. <https://doi.org/10.1016/j.caeo.2026.100399>
- Böhmer, A., Dillig, M., Andronache, D.-C., Beshlei, O., Bolaji, S., Isso, I., Kovaliuk, Y., Mason, J., Laza Medina, A., Stănescu, M.-H., & Yaroshenko, O. (2026). Conceptualizations of GenAI and students’ professionalization: Within the multi-layered environment of learning for higher education. *Computers and Education: Artificial Intelligence*, 11 (2026), 100636. <https://doi.org/10.1016/j.caeai.2026.100636>
- Carolus, A., Koch, M. J., Straka, S., Latoschik, M. E., & Wienrich, C. (2023). MAILS – Meta AI literacy scale: Development and testing of an AI literacy questionnaire based on well-founded competency models and psychological change- and meta-competencies. *Computers in Human Behavior: Artificial Humans*, 1(2), 100014. <https://doi.org/10.1016/j.chbah.2023.100014>
- Chee, H., Ahn, S., & Lee, J. (2025). A competency framework for AI literacy: Variations by different learner groups and an implied learning pathway. *British Journal of Educational Technology*, 56(5), 2146–2182. <https://doi.org/10.1111/bjet.13556>
- Chiu, T. K. F. (2025). AI literacy and competency: definitions, frameworks, development and future research directions. *Interactive Learning Environments*, 33(5), 3225–3229. <https://doi.org/10.1080/10494820.2025.2514372>

- Chiu, T. K. F., & Rospigliosi, P. "asher." (2026). Six global frameworks for human-centred AI literacy and competency: comparative analysis and a way forward. *Interactive Learning Environments*, 34(3), 1003–1005. <https://doi.org/10.1080/10494820.2026.2648342>
- Fan, Y., Tang, L., Le, H., Shen, K., Tan, S., Zhao, Y., Shen, Y., Li, X., & Gašević, D. (2025). Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance. *British Journal of Educational Technology*, 56(2), 489–530. <https://doi.org/10.1111/bjet.13544>
- Fu, Y., & Weng, Z. (2024). Navigating the ethical terrain of AI in education: A systematic review on framing responsible human-centered AI practices. *Computers and Education: Artificial Intelligence*, 7, Article 100306. <https://doi.org/10.1016/j.caeai.2024.100306>
- Isaeva, R., Caner, H. N., Caner, M., Giray, L., & Karadag, E. (2026). Students' engagement with generative AI in academic learning: A self-determination theory and epistemic network analysis study. *Computers and Education: Artificial Intelligence*, 10 (2026), Article 100606. <https://doi.org/10.1016/j.caeai.2026.100606>
- Khosravi, H., Shum, S. B., Chen, G., Conati, C., Tsai, Y.-S., Kay, J., Knight, S., Martinez-Maldonado, R., Sadiq, S., & Gašević, D. (2022). Explainable Artificial Intelligence in education. *Computers and Education: Artificial Intelligence*, 3, Article 100074. <https://doi.org/10.1016/j.caeai.2022.100074>
- Ma, M., Ng, D. T. K., Liu, Z., & Wong, G. K. W. (2025). Fostering responsible AI literacy: A systematic review of K-12 AI ethics education. *Computers and Education: Artificial Intelligence*, 8, Article 100422. <https://doi.org/10.1016/j.caeai.2025.100422>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, Article 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- OECD, & European Commission. (2026). Empowering learners for the age of AI: An AI literacy framework for primary and secondary education. OECD Publishing. <https://doi.org/10.1787/65cd27d4-en>
- Ouaazki, A., Shibani, A., Knight, S., & Holzer, A. (2026). Generative AI-enhanced learning experiences for computational thinking: A systematic scoping review and design guidelines. *Computers and Education: Artificial Intelligence*, 10 (2026), Article 100608. <https://doi.org/10.1016/j.caeai.2026.100608>
- Palod, V., Biswas, U., & Kambhampati, S. (2026). Evaluating the false trust engendered by LLM explanations. arXiv:2605.10930.

- Schorr, I., Bardach, L., Bühler, B., & Kasneci, E. (2026). Learning-to-learn in the age of generative AI: A scoping review and conceptual framework. *Computers and Education: Artificial Intelligence*, 10 (2026), Article 100575.
<https://doi.org/10.1016/j.caeai.2026.100575>
- Simms, B. (2025, November 29). Multimodal AI prompting for government applications. Graduate School USA. <https://www.graduateschool.edu/learn/ai/multimodal-prompting-in-government-ai-applications>
- Sofkova Hashemi, S. (2026). Students' multimodal prompting practices as epistemic work in AI literacy development. *Computers and Education: Artificial Intelligence*, 11 (2026), 100635. <https://doi.org/10.1016/j.caeai.2026.100635>
- Tankelevitch, L., Kewenig, V., Simkute, A., Scott, A. E., Sarkar, A., Sellen, A., & Rintel, S. (2024). The metacognitive demands and opportunities of generative AI. In *Proceedings of the CHI Conference on Human Factors in Computing Systems (CHI '24)*. ACM.
<https://doi.org/10.1145/3613904.3642902>
- U.S. Department of Labor, Employment and Training Administration. (2026, February 13). Artificial intelligence literacy framework (Training and Employment Notice No. 07-25).
<https://www.dol.gov/agencies/eta/advisories/ten-07-25>
- Xu, W., Dainoff, M. J., Ge, L., & Gao, Z. (2023). Transitioning to human interaction with AI systems: New challenges and opportunities for HCI professionals to enable human-centered AI. *International Journal of Human–Computer Interaction*, 39(3), 494–518.
<https://doi.org/10.1080/10447318.2022.2041900>
- Zhang, S., Xiao, R., Botelho, A. F., Liao, G., Chiu, T. K. F., Stamper, J., & Koedinger, K. R. (2026). How to assess AI literacy: Misalignment between self-reported and objective-based measures. *arXiv*. <https://arxiv.org/abs/2601.06101>

FIELD COMMENTARY (CITED AS PERSPECTIVE, NOT PEER-REVIEWED EVIDENCE)

The Chapter 3 section “A Fifth Consideration” draws on three interviews, not empirical research: cited at that lower evidentiary weight throughout, consistent with this ebook’s rule that every claim is stated at the strength its source actually supports. All three are interviews conducted and published by Allison Dulin Salisbury for *The Humanist* (Substack); the interviewee, not Salisbury, is the source of the view being cited, so in-text citations name the interviewee with “in Salisbury” to keep both facts visible. The Introduction’s “A

Personal Connection” section draws on the same lower evidentiary tier: Amarda Shehu, a George Mason University colleague, writing in her own voice on her own Substack (Shehu, 2026), not a peer-reviewed source, and not represented as one.

Salisbury, A. D. (2025, December 15). Can AI be pro-worker? A former labor secretary's perspective [Interview with S. Harris]. *The Humanist* (Substack).

<https://humanistxyz.substack.com/p/can-ai-be-pro-worker-a-former-labor>

Salisbury, A. D. (2025, December 22). George Siemens on why universities need to change [Interview with G. Siemens]. *The Humanist* (Substack).

<https://humanistxyz.substack.com/p/george-siemens-on-why-universities>

Salisbury, A. D. (2026, July 20). We're raising the first AI generation. They're pushing back [Interview with R. Winthrop]. *The Humanist* (Substack).

<https://humanistxyz.substack.com/p/were-raising-the-first-ai-generation>

Shehu, A. (2026, June 14). What do students want of AI? *Amarda Shehu* (Substack).

<https://amardashehu.substack.com/p/what-do-students-want-of-ai>

AI4LD eBook Series: Grounded in How People Learn. Designed for AI-Enabled Learning Experiences.

Updated July 2026 | Dr. Brenda Bannan | ai4ld.com