

AI4LD · BUNDLE C · UNIT 1

APPLIED

AI Literacy Curriculum Design Beyond Prompt Engineering

Why the AI literacy module your organization is asking for needs four dimensions, not prompting tips — and how to design it from the evidence.

Brenda Bannan, PhD

Professor Emerita, George Mason University · Founder, AI4LD

AI4LD · ILLUMINATING THE CONNECTION BETWEEN LEARNING AND AI · EBOOK + VIDEO
+ AI DESIGN COMPANION

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. We will not always get it perfectly right — and 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.

[AI Literacy Curriculum Auditor →](#)

A note on what the Auditor does and does not do: it scores coverage against a framework, drawing on the design logic laid out in this ebook. 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 — and 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 — the capacity to produce a desired output from a generative AI tool through instruction, iteration, and refinement. It is real, it is learnable, and it matters. 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 academic hair-splitting is direct: 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

Almost every instructional designer and LXD I talk with who has been handed “build an AI literacy module” as a brief describes the same instinct: start with the tool, start with the prompt, start with what is demonstrable in a single session. It is not a bad instinct. It produces something learners can practice and something stakeholders can see. But it also produces the same module, over and over, regardless of who is asking — because the tool is the easiest thing to teach, not the most complete thing to teach.

What the research reviewed in this ebook offers is not a reason to abandon that starting point, but a structure for what has to sit alongside it. The four-dimension framework gives you a way to check your own module against something more rigorous than “does this cover the tool” — and the evidence on where spontaneous practice consistently falls short gives you a specific, evidence-grounded answer to “what am I probably missing.” In your own practice, it is a continual struggle to keep up with the latest AI capabilities and releases, but some things hopefully remain unchanged such as honoring a human-centered design orientation in developing AI literacy.

What You’ll Learn

In the chapters ahead, we’ll cover:

The AI Literacy Evidence Base: What recent education research actually finds about how AI literacy develops — where it develops well on its own, where it doesn’t, and the specific failure mode (metacognitive laziness) that curriculum design has to counter.

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 — and 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 an outline. This chapter does that work for you. It reviews four findings from recent education research — the uneven four-dimension structure of AI literacy, prompting as an evolving practice rather than a one-time skill, the gap between using AI and understanding AI, and the specific failure mode curriculum design has to guard against — before Chapter 2 turns them into a design framework.

AI Literacy Is Not a Single Skill — It's Four Uneven Dimensions

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 the framework proposed in an exploratory review of thirty peer-reviewed studies (Ng, Leung, Chu, & Qiao, 2021) and operationalized directly in the education research this ebook draws on (Sofkova Hashemi, 2026).

The reason this structure matters for curriculum design is what happens when you apply it to real student practice: it does not develop evenly. In an exploratory workshop study analyzing 28 postgraduate students' self-documented prompting histories against this same four-dimension framework, students' reflections aligned closely with three of the four dimensions — a growing understanding of how GenAI interprets prompts, increasingly strategic crafting and refining of them, and increasingly critical evaluation of outputs for creativity, coherence, bias, and alignment with their intentions (Sofkova Hashemi, 2026). The ethical and human-centered dimension is where the pattern differed: students showed some awareness of bias and representational issues — dissatisfaction with stereotypical image representations led them to refine their prompts,

which the study reads as growing capacity to critically evaluate and create with AI — but the ethical and human-centered dimension itself “remained less explicitly developed” (Sofkova Hashemi, 2026). The encounters were productive; the growth they produced simply registered in the evaluative and creative dimension, not as explicit ethical reasoning. That is a single exploratory study, not a settled law — but its shape is exactly the design problem this ebook addresses.

Design implication: if you are designing a single module and have to prioritize, prioritizing the dimension least likely to develop without deliberate instruction is a defensible design decision — and the evidence points specifically to the ethical/human-centered dimension as that dimension.

Prompting as Epistemic Work, Not a Technical Instruction

A second finding reframes what “teaching prompting” should even mean. 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, dialogic 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 students learn to interpret, negotiate, and critically evaluate AI-generated content, not a technical instruction they either know or don’t.

Students in the same study also discovered something about the system’s behavior that no prompting cheat-sheet would have told them: a **structural rigidity** in outputs, where the model tended to preserve the underlying structure of an early draft — its “skeleton” — even under substantial revision requests, changing surface details more readily than deep structure (Sofkova Hashemi, 2026). That is a finding about system behavior that students had to discover through iterative practice, not a rule a one-time prompting lesson could hand them in advance.

Design implication: a prompting component built as a single early lesson, rather than as iterative practice sustained across a module, will under-teach even the “using” dimension it is meant to cover. Epistemic work requires repeated cycles of attempt,

observe, adjust — not a one-time demonstration.

Using AI Well Doesn't Mean Understanding AI

A separate line of evidence complicates the assumption that operational fluency produces conceptual understanding as a byproduct. In a study of 67 pre-service teachers practicing intercultural conflict scenarios with an AI-powered roleplay tool, the researchers found that 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 interacted with the same tool in a similar way (Böhmer et al., 2026). More pointedly: across the sample, using the tool effectively (technique/tool-oriented engagement) and understanding how or why it worked (structural engagement) were only weakly correlated — the authors' own framing is that **functional engagement with the tool does not automatically produce structural understanding of it** (Böhmer et al., 2026).

The same study reframes the AI system's instructional role through Vygotsky's Zone of Proximal Development: a generative AI tool can partially occupy the role of a "More Knowledgeable Other," but the authors are explicit that this mediation is technologically derived rather than genuinely social — "non sui generis" — and that this creates an open question about who bears responsibility for which parts of the learning process as the instructor's own role shifts toward what the authors call an "accompanying prompt engineer" (Böhmer et al., 2026).

Design implication: operational practice — however extensive — is not a substitute for direct instruction on how generative AI systems actually function. The two have to be designed for separately, because the evidence shows they do not reliably transfer from one to the other.

The Failure Mode This Curriculum Design Must Counter

The fourth foundational finding names the specific risk that makes deliberate curriculum design necessary rather than optional. A scoping review of 38 studies on GenAI-supported computational thinking instruction found a recurring pattern the field has named **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 — 66% of studies reported improved learning outcomes, but 16% reported decreases and 16% reported no change — which means “AI-enhanced” is not, on its own, a guarantee of better learning; design choices determine which outcome a given module produces.

A companion finding from a broader scoping review of learning-to-learn competencies gives this risk a name from outside the AI literature: the well-documented “**Google effect**” on memory, where reliable external availability of information leads people to retain knowledge of *where* to find something rather than the information itself, is offered as the direct analogue for GenAI overreliance — exposure to a GenAI tool does not, on its own, build the self-regulated learning it is often assumed to support, and excessive delegation produces what the authors call an “illusion of competence” (Schorr et al., 2026).

Notably, the risk runs in both directions. The same computational-thinking review found overreliance concentrated among beginners, while more advanced learners showed the opposite problem — under-utilizing GenAI tools on complex tasks where the tools could have genuinely extended 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 it cannot use a single design for learners at different skill levels. Guarding against metacognitive laziness for beginners and guarding against under-utilization for advanced learners are two different design problems, and a module needs to know which learner population it is designing for.

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-teaching, and vulnerable to a specific failure mode that only deliberate design can counter. Chapter 2 turns this evidence into a structure you can design against directly.

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

Using the Framework as a Design Lens

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

Dimension 1: Understanding AI

WHAT IT IS

Understanding AI is conceptual and structural literacy — knowing what a generative AI system is, broadly how it produces output, and what its genuine capabilities and constraints are. This 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

Learners who engage with AI tools more technically and structurally — using more tool- and-technique language, connecting their use to institutional and professional-practice domains — tend to be those with more advanced prior context; in the pre-service teacher study, this was the graduate cohort rather than the undergraduate cohort (Böhmer et al., 2026). But the same study found this was not a simple function of exposure: technique-oriented engagement and structural understanding were only weakly correlated across the sample as a whole.

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 (Böhmer et al., 2026). 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. Left to develop through use alone, this dimension is the one most likely to be shallow even in learners who otherwise seem fluent.

DESIGN IMPLICATIONS

Teach the mechanism directly — do not assume it will be inferred from practice.

Build a module component that explicitly addresses 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. Do not rely on hands-on practice alone to produce this understanding; the evidence indicates it frequently does not (Böhmer et al., 2026).

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, without reducing the time spent on hands-on practice.

Dimension 2: Using AI

WHAT IT IS

Using AI is operational literacy — 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.

WHAT THE EVIDENCE SHOWS

Operational skill develops through iterative practice, not a single instructional pass. Students' prompting strategies evolved from basic input-output requests toward more sophisticated, negotiated approaches — including asking the system to articulate criteria before applying them, and using role-play framing to shape tone — over the course of a structured, multi-session task (Sofkova Hashemi, 2026). This same study is the direct evidence for the modality gap named below.

WHERE SPONTANEOUS PRACTICE FALLS SHORT

This is the **multimodal prompting gap**: text-prompting competence does not 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.

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-specification 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, rather than assuming the system will infer it (Sofkova Hashemi, 2026).

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.

Dimension 3: Evaluating and Creating with AI

WHAT IT IS

Evaluating and creating with AI is the critical-application dimension — 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, it looked like a shift from basic interaction toward actively refining prompts to improve alignment with intentions, interrogating outputs, requesting explanations, and iteratively adjusting inputs to compensate for system limitations — including refining prompts when image outputs defaulted to stereotypical representations, a dissatisfaction-driven loop the study reads 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**: in a case study of 23 undergraduates, students consistently paired accounts of AI’s usefulness with active verification practices — corroborating AI output against course materials, requesting citations, cross-checking sources — and network analysis of their interview discourse showed this verification behavior was systematically, not coincidentally, linked 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 — 19 of 23 students described AI as an assisting tool, 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. The computational thinking scoping review found the opposite pattern present in a substantial share of the literature: **metacognitive laziness**, reduced critical engagement when learners over-rely on AI output, reported in 44% of the studies reviewed (Ouaazki et al., 2026). The same review found this risk concentrated among beginners specifically, while more advanced learners showed the inverse problem of under-utilizing AI on complex tasks — a tension the authors describe using the framing of a **"Zone of No Development"**: too much unstructured support collapses a novice's need to develop independent judgment, while too little contextualized support leaves an advanced learner's capability underused (Ouaazki et al., 2026).

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, is the more likely failure mode at that skill level (Ouaazki et al., 2026).

Practical example: A compliance training program gives all employees, regardless of tenure, the same unrestricted access to an AI research assistant during a case-analysis exercise. New hires produce plausible-looking analyses that go unverified against the actual policy documents; experienced staff, given the same tool with the same

instructions, barely use it because the unstructured access doesn't map onto how they already work through complex cases. Redesigning the exercise with two tracks — a scaffolded, verification-required track for new hires and an open-ended, tool-integrated track for experienced staff — improves both groups' outcomes 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 — 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 the evidence returns to as the one that does not develop explicitly through unstructured practice — stated carefully, at the strength the methods support. In the exploratory postgraduate workshop study, students did encounter ethical material directly — stereotyped defaults in AI-generated images (such as unspecified skin tone) and inconsistent, differently transparent copyright refusals across tools — and they engaged with what they encountered: dissatisfaction with stereotypical representations, particularly around authenticity and appropriateness, led them to refine their prompts, a progression the study interprets as growing capacity to critically evaluate and create with AI and as a critical understanding of the limits of representation in GenAI (Sofkova Hashemi, 2026). What the study reports about the ethical and human-centered dimension itself is more specific: students “showed some awareness of bias and representational issues,” but this dimension “remained less explicitly developed” than the other three (Sofkova Hashemi, 2026). The encounters were not wasted — their yield simply landed in evaluation and creation, not in explicit ethical reasoning about what should be made, shipped, or shown.

The strongest version of this finding comes from outside the prompting literature entirely. A scoping 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 directly still: **the authors found no studies examining GenAI's use to actively teach the ethical dimension of learning-to-learn at all** — a stated blind spot in the field itself, not a criticism of any single program or curriculum (Schorr et al., 2026).

A separate, independent line of evidence reaches the same design conclusion from a different literature. The computational thinking scoping review's seven design guidelines include, as its own distinct guideline, designing with digital ethics and integrity in mind — bias, privacy, transparency — named separately from the other six guidelines specifically because the review found it was not otherwise reliably addressed by existing GenAI-in-education implementations (Ouaazki et al., 2026).

WHERE SPONTANEOUS PRACTICE FALLS SHORT

Everywhere the evidence led us, the sources seemed to fall short. To be precise about what kind of gap this is: it is not that the field lacks ethical frameworks or guidance — AI ethics scholarship, institutional principles, and dedicated research centers 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 education research now systematically reviewed in their own right (Fu & Weng, 2024). Nor is the ethical dimension untaught everywhere: 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 is missing is the translation step into adult and professional AI-supported learning — the contexts instructional designers and LXDs actually build for and evaluate. In the empirical GenAI literature reviewed here, that scholarship has not yet become tested instructional models that turn ethical guidance into taught, practiced, assessed reasoning: the scoping review found limited evidence of studies that used GenAI to actively teach the ethical dimension (Schorr et al., 2026). And the exploratory workshop evidence shows what happens in the absence of that design: encounters with bias and representational issues do real work — they drive prompt refinement and critical evaluation — but that growth registers in the evaluating-and-creating activities,

while the ethical and human-centered dimension is less explicitly developed (Sofkova Hashemi, 2026). One exploratory study cannot settle the question; however, it can, and does, illustrate the shape of the gap.

DESIGN IMPLICATIONS

Design this dimension as its own explicit unit — do not assume it will emerge from encounters with the tool. Because the empirical literature reviewed here offers no validated pedagogy for this dimension in adult and professional contexts, the safest design approach is a conservative one: structured reflection, scenario-based discussion, and explicit naming of bias and integrity issues when they occur — drawing on the ethical frameworks that do exist, such as 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 that turns “I fixed the biased image” into “why did the system default there, and what does that mean for what we ship?”

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) — which suggests ethical instruction in an AI literacy module benefits from being explicitly framed as a human-guided conversation, not a feature of the AI tool itself.

Practical example: A university’s AI literacy workshop for staff includes a hands-on image-generation exercise. Several participants notice the tool defaults to a narrow range of skin tones and body types unless explicitly prompted otherwise — and discuss it informally among themselves during the break, but the observation never surfaces in the structured session. Adding a five-minute structured debrief immediately after the

exercise — “What did the tool default to, and what did you have to specify to change that?” — turns an incidental observation several participants already made into the module’s most substantive ethical discussion, at minimal additional design cost.

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)	Design as its own explicit unit; 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 prompting alone:

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.

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 when 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. 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 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: for beginners, favor tutor-style prompting patterns that require an attempt before AI assistance is available; for advanced learners, favor 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 two-track design decision, not a single setting. A module built for a mixed-experience audience needs to either branch by skill level or make an explicit, instructor-facing decision about which failure mode it is optimizing against for that specific audience.

Design Strategy 4: Design the Ethical Dimension as Its Own Unit

This is the strategy the evidence most consistently supports as necessary and most consistently shows is missing. No study in one literature review revealed how to actively teach the ethical dimension 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).

Practical design approach: Treat this as a designed unit with its own objective, not a hoped-for byproduct of hands-on practice. Build structured reflection prompts around the ethical encounters your other exercises will likely already create — biased defaults, inconsistent policy behavior, confidently wrong output. Because no validated pedagogy seems to exist yet for this dimension specifically, and it currently seems to favor conservative, discussion- and scenario-based methods over any single supported technique.

Example: A healthcare organization's AI literacy module for administrative staff includes a hands-on exercise generating patient-facing communications. The design team adds a short structured debrief immediately after the exercise, asking staff to identify any assumption the tool made by default that they had to correct — about tone, formality, or audience — and to discuss, as a group, what it would mean for that same default to go uncorrected in a real patient communication. The debrief costs ten minutes and turns a technical exercise into the module's clearest example of ethical, human-centered design judgment.

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 — Rebecca Winthrop and 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) — outside the peer-reviewed literature this ebook otherwise draws on, and named here as field perspective rather than evidence — describes a real and reasonable skepticism by learners: concern about job security, about cognitive offloading, about whose interests a given AI deployment actually serves. Siemens in particular argues the deeper issue is institutions adopting AI without deciding what knowledge and skills actually stay stable — 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.

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?
Does the module have a distinct, designed component for the ethical/human-centered dimension?	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 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 — 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 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

- 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>
- 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>
- 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>

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>

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>

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>

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.

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>

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

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