

Teaching AI Responsibly to Principled Skeptics

A Framework for Students, Faculty, and Administrators

James M. Hyman

Department of Mathematics, Tulane University

`mhyman@tulane.edu`

2026-08-26

Abstract

Universities are being urged to make AI use universal, and students who decline on environmental, ethical, or intellectual grounds are treated as an obstacle to that goal. This paper argues the reverse. They are not the obstacle. They are the asset.

The argument rests on a claim developed here and called the formation precondition. The competencies that AI literacy frameworks specify are formed through unaided practice, and the tool those frameworks concern is the one most able to displace that practice. Protected periods of unaided work are therefore not a courtesy extended to objectors. They are a condition on which any framework of AI competencies depends, which makes principled refusal a design requirement rather than an exemption to grant, and makes it hold for students who raise no objection at all.

The paper sets out a five-layer account of AI literacy, technical, environmental, ethical, practical, and civic, in which technical literacy is the base the others rest on and environmental cost is treated as a literacy in its own right rather than an aside. It then gives separate guidance to three audiences: students deciding how to learn, faculty designing courses, and administrators setting policy. For faculty it shows how to teach the skeptic without either overriding the objection or excusing the student from the outcome. For administrators it separates a literacy obligation from a use mandate, a distinction the European Union's AI Act draws in Article 4 and a 2026 amendment sharpens.

Keywords: AI literacy; generative AI in higher education; principled refusal; cognitive offloading; environmental cost of AI; course design; institutional policy

Working draft, v1.7.0.

1 Universal Framing

1.1 The Teaching Problem

A student in an environmental science class refuses to use ChatGPT for an assigned literature review. She has read about water consumption in Chilean towns near new data centers, and she will not

contribute to the demand. Her instructor pauses. There are three plausible responses. Require her to complete the assignment as written. Exempt her and let her substitute a different task. Or take her objection seriously enough to redesign the question she has actually asked, which is no longer about a literature review but about how a thoughtful person decides which technologies deserve their participation. This scene plays out, in different forms, in classrooms across higher education, on faculty senate agendas, and in administrative offices where AI policy is being written.

This document argues that the third response is the only one consistent with what higher education exists to do. AI is not ethically neutral, and the environmental, labor, and political costs the student is pointing to are real. But refusing to understand AI is not the same as resisting its harms. A responsible education should help students study AI critically, use it sparingly when appropriate, reject it when the evidence warrants, and advocate for better systems. Principled skeptics, environmentally concerned students prominently among them, are a population higher education has the strongest pedagogical reason to develop into rigorous AI critics. They are not the obstacle. They are the asset.

The argument addresses three audiences. Students must decide how AI fits into their education and their conscience, and whether using it makes them complicit in harms they oppose. Faculty must design courses that respect principled refusal while preserving learning outcomes, and respond to a classroom split between enthusiastic users and committed abstainers. Administrators must set institutional policies that support faculty discretion without abandoning either the skeptical student or the enthusiastic one, and they must do so while contracts, vendor relationships, and procurement decisions are already in motion. Each audience faces a different version of the same question. The framework developed here gives them shared vocabulary.

The document is a framework paper, not a research study or a policy template. It draws on institutional reports from the International Energy Agency, the OECD, UNESCO, and EDUCAUSE, on recent scholarship in AI literacy and the environmental footprint of large language models, and on the experience of teaching across decades of curriculum change. It does not catalog every AI literacy framework now in circulation, nor does it offer a complete inventory of AI's harms. It offers a framework that engages principled skepticism without dismissing it and without surrendering to it, and it names the condition on which any framework of the kind depends.

The opening scene is a composite. It draws on conversations the author has had with students and colleagues across decades of teaching, and on the published cases discussed in Parts 1.2 and 2.3 (in particular Crawford (2021) and Hao (2025)), but it is not a transcript of any single exchange. The composite form is deliberate: it lets the scene generalize across the institutional contexts the document addresses, and it protects any individual student from being treated as representative of a population that contains real internal diversity. Readers seeking documented case studies in this literature should turn to Crawford, Hao, and the grounding documents behind the CCCC 2026 resolution and the AHA 2025 principles. The document's contribution is the framework, not the case material.

Part 1 locates the problem and characterizes the principled-skepticism population the document responds to. Part 2 presents the five-layer framework that organizes the rest of the document. Part 3 translates the framework into guidance for each audience. Part 4 closes by returning to the opening scene. Readers can read straight through or skip directly to the audience section that applies to them.

1.2 Principled Skepticism: Environmental and Beyond

Students who object to AI on environmental grounds are making a moral argument, not merely a technical complaint. They are asking whether their personal participation contributes to a system they believe is harmful. The question deserves engagement on its own terms, not deflection through reassurance.

The data make the concern hard to dismiss. The International Energy Agency projects that global data-center electricity consumption could roughly double by 2030, reaching approximately three percent of global electricity demand (International Energy Agency, 2025). AI-focused data centers are expected to drive that growth at a faster rate than data centers overall.

The footprint extends beyond electricity. Experts convened by the OECD have described significant water demand for data-center cooling and for the electricity generation that supports it, with particular concern in water-stressed regions (OECD, 2025). Recent reporting has tracked specific cases of community-level resistance, including sustained community opposition to data-center water use in Chile (Hao, 2025). The environmental cost is local as well as aggregate.

A clear sense of scale matters. The environmental impact of a single text prompt is not the impact of building a data center, training a frontier model, or deploying AI across millions of commercial interactions. Google reported in 2025 that a median Gemini Apps text prompt used about 0.24 watt-hours of electricity, about 0.03 grams of CO₂ equivalent, and about 0.26 milliliters of water (Google Cloud, 2025). The figure warrants caution because it comes from a company with an interest in presenting AI inference as efficient, but it is useful for distinguishing personal-use concern from infrastructure-level analysis. Independent estimates of per-prompt energy and water use vary substantially from the Google figure depending on model size, hardware, and infrastructure assumptions, with the Google figure sitting toward the low end of the published range. The order-of-magnitude variation does not change the personal-versus-infrastructure distinction, but it should temper any specific number's authority. A student may believe that using AI once for an assignment is morally comparable to endorsing the entire industry, when the two questions are usefully separated.

The strongest student objection rarely rests on energy alone. Many activities use power. The deeper concern is that AI generates new energy demand while also concentrating economic power in a small number of companies, weakening labor protections, enabling surveillance, amplifying misinformation, and undermining ownership of creative work. Karen Hao's investigative account of OpenAI documents specific instances of each, including the exploitation of data-labeling workers in Kenya alongside the water and electricity costs (Hao, 2025). Kate Crawford's earlier framing of AI as a planetary-scale extractive industry, drawing minerals, energy, water, and labor from communities that do not share in the resulting wealth, gives the same observation broader theoretical reach (Crawford, 2021).

The document develops the environmental case at the greatest depth because environmental concern is the entry point through which many students arrive at principled AI skepticism, and because environmental scale (energy, water, carbon) has documented infrastructure-level evidence that the other concerns do not yet have to the same extent. The labor, privacy, surveillance, misinformation, and ownership concerns operate alongside the environmental one and warrant comparable framework development in further work. This document treats them as illustrative of the principled-skepticism population's range, not as fully developed parallel cases. A reader whose lead concern is labor or privacy rather than environmental will find the framework's five layers translate, but should expect the worked examples and assignment templates to require adaptation.

The concern has been documented in the technical literature for years. Strubell et al. (2019) quantified the financial and environmental costs of training large language models in 2019. The research conversation on those costs has run since. Bender et al. (2021), in the paper that prompted Timnit Gebru's departure from Google, argued that the environmental and financial costs of large language models should be weighed before development rather than after. The technical discourse on AI's environmental costs has been serious within the field for more than half a decade. Some of that discourse has reached students through downstream coverage in mainstream and student media; the students raising the concern in classrooms are responding to that diffuse but related body of public discussion rather than to the technical literature directly.

A useful distinction sits in the middle of this argument. Consumer refusal is the choice to abstain from casual use, comparable in form to avoiding disposable plastic or fast fashion. Civic literacy is the capacity to evaluate and challenge a technology that is already embedded in institutions, employers, governments, schools, health systems, and media platforms. The two are not in conflict. A student may refuse personal AI use and still benefit from understanding how the systems work well enough to argue against their harmful applications. Refusal without literacy is moral expression. Literacy without refusal is technical competence. The combination, refusal joined to literacy, is what democratic AI governance requires.

The refusal position now has formal institutional backing. The Conference on College Composition and Communication, the largest professional organization of writing educators, passed a resolution in March 2026 affirming the right of students and faculty to refuse the use of generative AI in the writing classroom (Conference on College Composition and Communication, 2026). The American Historical Association's guiding principles acknowledge that some educators have rejected generative AI for its ethical, environmental, and economic consequences, while arguing that ignoring the technology will neither halt its spread nor shield the discipline from its reach (American Historical Association, 2025). The skeptical student is not isolated. She is part of a movement that includes her professional society leadership.

1.3 Stakes Across Roles

The same teaching problem has different stakes for different audiences. Each version of the problem requires different responses, but the cost of getting it wrong points in the same direction: AI literacy will be developed somewhere, by someone, on some terms. The question is whose terms.

For students, AI literacy is becoming a functional competency for professional and civic life. The European Union's AI Act, in Article 4, establishes a formal requirement for AI literacy across organizations that develop and deploy AI systems (European Union, 2024), a signal that civic and professional AI literacy will increasingly carry regulatory weight. A student who declines to engage with AI on principled grounds may still face an employment market, a credentialing landscape, and a civic environment in which the capacity to evaluate AI is assumed. Refusing to use AI is a defensible personal stance. Refusing to understand it leaves the student less equipped to challenge the systems they oppose.

For faculty, the stake is whether the classroom remains a place where students can develop their own positions on contested technologies, or becomes a space where AI is either uncritically adopted or quietly forbidden. EDUCAUSE's 2024 framework for AI literacy in teaching and learning argues that students, faculty, and staff need technical, evaluative, practical, and ethical understanding

of AI in academic and professional contexts (EDUCAUSE, 2024). A faculty member who treats AI as inevitable surrenders the curriculum’s critical function. A faculty member who treats it as forbidden leaves students unprepared for institutional environments where AI is already deployed. Both responses underserve the skeptical student, who came to higher education expecting better than either.

For administrators, the stake is whether AI policy at the institution is set deliberately or by default. Vendor contracts are being signed. Procurement decisions are being made. Faculty are forming their own classroom policies in the absence of clear institutional guidance. The Conference on College Composition and Communication resolution in March 2026 affirming the right to refuse AI in writing classrooms places a question on every administrator’s desk: how does the institution support faculty who refuse, faculty who teach with AI, and students at each end of that spectrum, simultaneously? Avoiding the question does not make it go away. It transfers the answer to vendors, to the most visible faculty incidents, and to the regulatory environment, whose AI requirements (the EU AI Act, increasingly state-level US legislation) are arriving regardless of institutional readiness.

The loudest counterpart to the position developed here is the workforce-preparation framing of AI literacy, advanced by industry-funded initiatives (the Microsoft AI Skills Initiative, comparable programs at large workforce-oriented institutions, and a growing class of vendor-sponsored curriculum partnerships). On that framing, AI literacy is what employers will require, and higher education’s job is to deliver it. The framing has the advantage of taking the technology seriously and the disadvantage of treating students as inputs to a labor market rather than as citizens with judgment to develop. The framework here does not reject the employability concern; the “For students” treatment above acknowledges that AI competency is becoming functional for professional life, and the EU AI Act’s Article 4 requirement is cited as evidence. What the framework rejects is letting employability set the syllabus. A student who learns AI on workforce-preparation terms learns to use it; a student who learns AI on civic-literacy terms learns when to use it, when to refuse it, and how to argue for changing it. The first competency is contained within the second. The reverse is not true.

The remainder of this document treats the three audiences as sharing one problem but holding different responsibilities for it. The framework in Part 2 gives them shared vocabulary. The audience-specific guidance in Part 3 gives each of them concrete work.

2 The Framework

2.1 A Framework for AI Literacy

Five forms of literacy organize the rest of this document. They are introduced together because they reinforce each other. None of them is sufficient on its own; all five become available to a student when the others are also present. The five are technical, environmental, ethical, practical, and civic. Each is described below in terms of what understanding it requires, what teaching it looks like in concrete classroom moves, and what a student who has developed it should be able to do.

The order of the five is not arbitrary. Technical literacy makes the other four possible: a student who does not understand what a model is doing cannot weigh its environmental costs, identify its ethical stakes, use it well, or argue for better policy. Environmental literacy follows because environmental

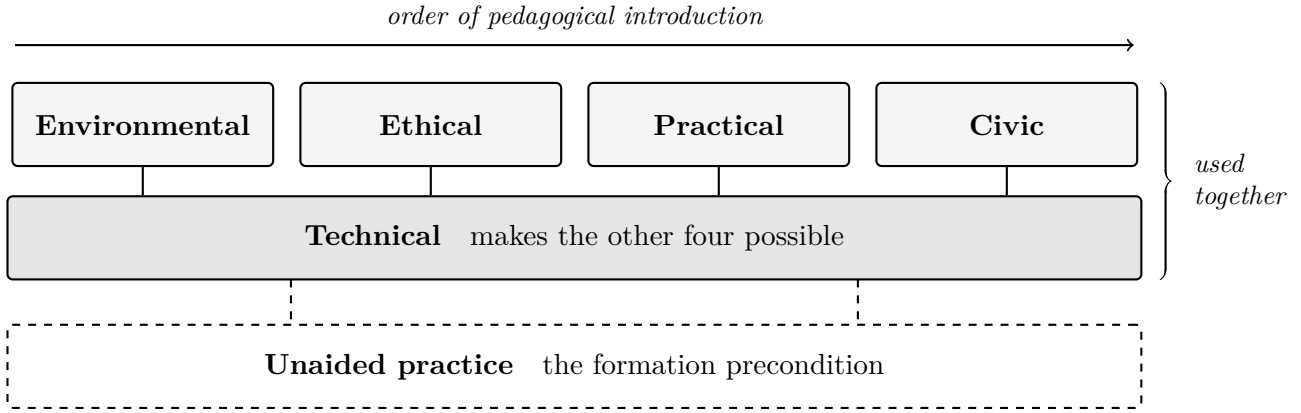


Figure 1: The five literacies and the condition beneath them. Technical literacy is the base the other four rest on, and environmental, ethical, practical, and civic are introduced in that order while being drawn on together whenever a student evaluates a real application. *What to notice:* the dashed foundation is not a sixth literacy. It is the unaided practice through which all five are formed, and it is the one element here that a course can remove without intending to.

concern is the entry point through which many students arrive at AI as a serious subject. Ethical literacy widens the lens. Practical literacy converts understanding into capability. Civic literacy returns the student to the public questions that brought them into the classroom. The five are sequential in their pedagogical introduction but simultaneous in their use: a student evaluating a specific AI application is drawing on all five at once. Figure 1 shows the arrangement, together with the condition developed in Section 2.2. Frameworks of this shape are now numerous. Long and Magerko (2020) set out seventeen competencies. Ng et al. (2021) reduced them to four aspects. Later schemes propose five, seven, and eight dimensions for higher education and for teacher preparation. Little is gained by adding another count, and this framework does not claim its contribution lies there. Two things distinguish it. The first is that environmental cost is treated as a literacy with its own competencies rather than as one item inside a broader social or ethical category. The second is the subject of Section 2.2: the framework states the condition on which it, and every framework built this way, depends.

Technical literacy means understanding what generative AI does at a level of detail sufficient to reason about its outputs. Students do not need the mathematics of neural networks. They need to understand that a large language model is a system trained to predict likely text completions from patterns in training data, that “knowing” something for such a system means having seen it represented frequently in a particular form, that confident-sounding output can be confidently wrong, and that the gap between fluency and reliability is the central thing a human reader has to manage. A student with technical literacy can answer questions like: Why might a model produce a plausible-sounding but incorrect citation? What kinds of errors should I expect, and what kinds should surprise me? When does the model’s fluency mislead me about its accuracy?

Environmental literacy means understanding the scale distinctions developed in Part 1: between training and inference, between text and image generation, between individual use and industrial deployment, between average estimates and worst-case infrastructure effects. It also means understanding how environmental figures are produced, which figures are reliable, and which depend on assumptions that may not hold. A student with environmental literacy can read a company sustainability report and identify what is being measured, what is being omitted, and what comparison

would make the figure interpretable. They can distinguish a personal-use question (Is this prompt wasteful?) from a structural question (Is the deployment of this technology consistent with the institution’s stated environmental commitments?).

Ethical literacy covers the catalog of concerns that AI ethicists have organized: privacy of training data, bias in model outputs, labor conditions in AI development, copyright and creative ownership, misinformation risk, academic integrity, surveillance applications, and accountability gaps when a system causes harm. The UNESCO ethics frameworks and the OECD AI Principles provide useful entry points (UNESCO, 2021; OECD, 2024), but ethical literacy is more than familiarity with frameworks. It is the capacity to identify which ethical concerns are at stake in a specific use case and to weigh them against each other when they conflict. A student with ethical literacy, asked whether to use AI for an assignment, can name three or four distinct ethical concerns that bear on the choice and explain why one might outweigh another in the specific context.

Practical literacy is built through use, not assertion. Students need experience with AI that produces plausible errors, that wastes their time, that weakens their thinking on a task they had been doing well, and with AI that helps them see something they had missed. The literacy is the discrimination between those situations. A student with practical literacy can predict, before submitting a prompt, what kind of output they are likely to get, whether that output will be useful, and how much verification effort the use will require to be worth doing. The assignments in Part 3’s faculty section are designed to build this discrimination. The practical literacy of a critic is sharper than the practical literacy of an enthusiast, because the critic has spent more time examining what is wrong with the output.

Civic literacy connects individual and institutional AI practice to the policy questions that surround them. The EU AI Act’s literacy requirement under Article 4 provides a concrete regulatory context for why this matters (European Union, 2024). Questions about data-center transparency, environmental reporting standards, procurement requirements, vendor accountability, and public oversight are not abstract: they are the levers through which AI’s harms can be reduced. A student with civic literacy can answer questions like: What does the institution disclose about its AI vendor contracts? What environmental reporting requirements apply to data centers in this state, and what would stronger requirements look like? Which of these questions can be answered by reading existing public documents, and which require advocacy to be answerable at all?

A course organized around these five forms of literacy is not promoting AI. It is preparing students to live and work responsibly alongside it, with the technical understanding to evaluate its outputs, the environmental understanding to weigh its costs, the ethical understanding to identify what is at stake, the practical understanding to use it well or refuse it well, and the civic understanding to advocate for systems that would govern it better. The skeptical student arrives in the classroom with strong instincts in several of these layers already. The pedagogical task is to develop those instincts into rigorous competence, not to redirect them.

2.2 The Condition the Framework Assumes

Every literacy described above is built through practice. That is easy to say and easy to read past, so it is worth stating as a condition rather than leaving it as an aside.

Consider what practical literacy requires. A student learns to tell useful output from plausible error by having been reliable without help first. The description above says the literacy is built partly

through experience with AI that weakens their thinking on a task they had been doing well. That sentence assumes a task the student had been doing well. Remove the assumption and nothing remains to notice. There is no felt difference between working with help and working without it, because there is no memory of working without it.

The same holds for technical literacy, which the framework treats as the base the other four rest on. Technical literacy is the capacity to manage the gap between fluency and reliability. A student manages that gap by holding an independent standard against which confident output can be checked. That standard does not come from reading about it. It comes from doing the work unaided often enough to know what a sound answer feels like.

So the framework rests on something it has not named. It assumes students who still do enough unaided work to build the discriminations it describes.

The assumption is not safe. Work on cognitive offloading has documented what happens when it fails. In a randomized study of writing tasks, students working with generative AI raised their task scores while monitoring and evaluating their own work less, and gained nothing in transfer to later problems (Fan et al., 2025). The fluency of the output creates an impression of competence the learner has not acquired. Fan et al. (2025) call the pattern metacognitive laziness.

Two bodies of work grew up apart. One asks what an AI-literate person should know and answers with dimensions: four, five, seven, or eight, depending on the scheme (Hackl et al., 2026). The other asks what AI use does to learning and answers with evidence about displaced effort. They meet often now. The AI literacy framework published jointly by the OECD and the European Commission cites the metacognitive laziness finding by name, reports that AI-assisted work need not produce durable learning, and warns that students who lean on the tool are underprepared when it is taken away (OECD/European Union, 2026).

What the meeting has not changed is the frameworks. In that document the finding sits in a discussion of risk, and none of the four competence domains asks that any competence be formed unaided. The closing instruction, to prioritize valued human knowledge regardless of AI use, names the destination and leaves the road unmentioned. Where the meeting does reach the framework itself, it produces developmental models: stages a learner passes through, with uncritical use treated as an early condition to be outgrown. Liu and Levy (2026) give the fullest version. Their first stage holds both principled refusal and simple inexperience, and the destination they recommend for graduates lies several stages further on.

That treats non-use as a starting point. The argument here is that it is also a condition. The competencies a dimensional framework specifies are acquired through practice, and the tool the framework is about is the thing that can displace the practice. A framework that does not say so is not simply incomplete. It describes a destination while the road is being taken up.

The disagreement is narrow and worth stating exactly. Liu and Levy (2026) are careful to say that institutions should be clear about what options exist for students who want to minimize their AI use, and that non-engagement is not a deficiency. The question is whether that provision is a courtesy offered at the start or a condition maintained throughout.

Call this the formation precondition: the competencies a literacy framework specifies are formed through unaided practice, and the tool the framework concerns is the one most able to displace that practice. Naming the condition changes what follows from it.

If unaided practice is a precondition on every literacy in this framework, then protected non-use

is not a courtesy extended to students who object. It is the mechanism that keeps the framework working. A course that requires AI use at every stage removes the ground its own learning outcomes stand on. A course that protects periods of unaided work is not accommodating the skeptical student. It is preserving the condition under which any student becomes AI literate at all.

That reverses the usual order of the argument. Principled refusal is normally treated as a position to be respected. Here it is a design requirement, and it holds for students who raise no objection whatever.

The consequence for course design is developed in Part 3. Stated briefly: an assignment meant to build AI literacy has to be sequenced so the formative work comes before or beside the assisted work, never in place of it. A student who has never solved the problem learns little from watching a model solve it.

The condition also marks the framework’s outer limit. Where students arrive having done almost no unaided work in a subject, the five literacies are not available to be taught in the order given, and the first task is not AI literacy at all.

2.3 Methodology and Source Assessment

This document draws on three kinds of sources. The first is institutional reports from organizations with broad responsibility for AI policy: the International Energy Agency, the OECD, UNESCO, EDUCAUSE, and the regulatory text of the EU AI Act. These provide the policy and governance framing. The second is peer-reviewed scholarship on AI literacy and on the environmental and ethical costs of large-scale AI systems, including the foundational work of Long and Magerko (2020), Ng et al. (2021), Strubell et al. (2019), and Bender et al. (2021). The third is contemporary critical reporting and book-length analysis, including Crawford’s *Atlas of AI* (Crawford, 2021) and Hao’s *Empire of AI* (Hao, 2025), which translate the institutional and scholarly findings into specific cases and accessible argument.

The institutional sources are strong for broad policy and educational framing. The peer-reviewed sources establish that the environmental and ethical concerns the document treats seriously have been treated seriously in the technical literature for years. The Google figure for single-prompt energy use, cited in Part 1, is useful for establishing scale but requires the caution noted there: it comes from a company with an interest in presenting AI inference as efficient. It should not be used to dismiss infrastructure-level concerns.

Three conclusions follow from the evidence with high confidence. AI’s aggregate energy and water demands are serious and growing. Single-prompt environmental concern is often less analytically useful than system-level analysis. AI literacy carries genuine ethical weight because AI is already deployed in consequential institutional settings, and the public’s capacity to evaluate it has not kept pace with its deployment. A fourth conclusion follows with moderate confidence: careful educational use of AI can be ethically justified when it teaches students to evaluate, limit, challenge, and govern AI use, provided that use is deliberate, documented, and subject to critical reflection.

A note on this document’s own production is warranted, both as transparency and as a worked example of the framework applied reflexively. This document was developed with AI assistance under a structured review protocol over the course of several drafting sessions. The framework, the argument, the judgments about what to include and what to leave out, and the words on the page

are the author's. The AI's role was to surface candidate framings, verify citations against primary sources, scan prose for the kind of signature stylistic patterns the framework's critical-observer mode would teach students to recognize, and stress-test the argument against the strongest objections a critical reader would raise. Applying the framework's own questions to this production: the use was deliberate, the AI's contribution is documented, the outputs were subject to critical reflection at each step, and the educational gain (a stronger argument, more rigorous citation practice) is weighed against the environmental and labor costs of the assistance itself. The document is not exempt from the framework. It is one of the framework's cases.

A line of critical scholarship on educational technology, represented by Selwyn (2024) and others, argues that even framing the response to AI as "literacy" can serve to normalize the technology and to accept its providers' framing of what literacy means. The argument is serious. The civic literacy component developed in Section 2.1 is intended in part as a response: a literacy that organizes students to ask whether AI providers should be granted their current latitude, rather than only how to use the tools they offer, is not the literacy AI providers would design. The framework here is offered in that spirit, not as a settled answer to the question the critical scholarship raises. Alternative organizing concepts (agency, accountability, technological citizenship) would face structurally similar co-optation pressures once they entered institutional channels; the question is which framing best resists co-optation rather than which framing avoids the risk entirely. "Literacy" survives the choice in this framework because it carries a built-in expectation that the learner develops independent judgment, not just operational competence, and because a literacy framework with civic literacy as an explicit dimension is among the framings best positioned to equip students to recognize when the framework itself could be turned against them. Readers persuaded that no institutional framework can resist co-optation may find the Selwyn position more defensible than the one developed here; the disagreement is real and the document does not claim to have closed it.

The framework has limits worth naming. It targets higher education and does not address K–12 contexts directly, though some adaptation is straightforward. Its institutional sources are Anglophone and reflect a US and European policy landscape; non-Western contexts have different regulatory environments and different histories with technology adoption that may shift the priorities. The civic literacy layer is the most context-dependent of the five: where regulatory environments do not give civic participation purchase on AI governance, or where institutional faculty governance does not exist in recognizable form, the civic literacy layer requires substantial adaptation rather than direct application. The empirical literature on student attitudes toward AI is still young: this document cites the most-referenced surveys (Chan and Hu, 2023, among others), but the population it characterizes is changing as the technology and its discourse evolve.

3 Audience-Specific Guidance

3.1 For Students Learning with AI

If you are a student reading this section, you may be anywhere along a wide spectrum. You may already use AI tools daily and want to use them better. You may use them reluctantly because your courses require it. You may refuse them entirely on environmental, ethical, or pedagogical grounds. You may use them for some tasks (brainstorming, summarization) but avoid them for others (drafting, decision-making). Wherever you stand, the framework in Part 2 asks the same

questions of you. The differences come in the answers.

Self-assessment using the five-layer framework

Technical literacy asks whether you understand what a language model actually does when you prompt it. If you cannot explain in your own words why a model might confidently produce a wrong citation, you are not yet technically literate in the sense this framework requires. The remedy is reading: a few clear introductions to how transformer-based models work will get you most of the way there.

Environmental literacy asks whether you can distinguish a personal-use question from an infrastructure question. The Chilean water cost is not directly affected by whether you submit one more prompt. It is affected by procurement decisions, deployment scale, and the policy environment your prompt sits within. If you find yourself agonizing over single prompts while paying no attention to your institution's vendor contracts, your environmental literacy is half-developed.

Ethical literacy asks whether you can name the specific ethical concerns at stake in a specific use case. "AI raises ethical concerns" is a slogan, not literacy. "Using AI to draft an essay raises concerns about my own learning, about the labor that produced the training data, and about my eventual ability to evaluate the output's quality" is closer.

Practical literacy asks whether you can predict, before you submit a prompt, whether the output will be useful. If you regularly find AI output disappointing, you are early in your practical literacy. If you predict that disappointment and prompt anyway, you are early in your practical wisdom.

Civic literacy asks whether you can identify which policy questions about AI you have the standing to act on. As a student, you can ask your institution what it discloses about AI vendor contracts. You can ask your senators what environmental reporting requirements they support for data centers. You can ask your professional society what positions it has taken. Civic literacy is the habit of asking those questions out loud.

A personal compact

A short personal compact, made to yourself rather than imposed on you, signals that your AI use is governed by your own values. Consider committing to the following:

1. Use AI only when it serves a clear learning purpose. Casual use is the kind that adds up environmentally and erodes the habit of thinking for yourself.
2. Prefer small text-based uses over energy-intensive media generation unless the task requires it.
3. Draft your own thinking before asking AI. The AI's answer is more useful to you when you already have one to compare it against.
4. Never submit AI output you do not understand. The signature on the work has to mean something, including to you.
5. Check important claims against reliable sources. Confident-sounding output is not the same as accurate output.

6. Do not enter private or sensitive information into commercial AI tools. Anything you submit becomes training data or risk.
7. Disclose AI use when required, and consider disclosing it even when not.
8. Ask, for any specific use, who benefits and who may be harmed.
9. Treat AI as a tool to examine, not an authority to obey.

If you choose not to use AI at all, the compact's logic still applies in reverse. The questions about purpose, harm, and accountability are the same. The answers are simpler.

How to use AI without being used by it

Prompt minimalism is the practice of doing more thinking before you prompt so that you prompt less. Iterating ten low-quality prompts costs more energy than one well-formed prompt. It also produces worse results, because each iteration adapts to the noise in the previous output. Before you prompt, decide what you actually want, what context the model needs, and what you will do with the output. A prompt that produces a usable answer the first time is almost always shorter, more specific, and preceded by more of your own thinking than one that doesn't.

Verification is the practice of treating AI output as a draft you must check, not an answer you can adopt. The model has no way of distinguishing between what it has been trained to predict and what is true. You do. When the output gives you a specific fact, a specific citation, or a specific recommendation, check it. When it gives you a structural suggestion, evaluate whether the structure fits your actual case.

Refusing specific uses is different from refusing all use. You may use AI for research help and refuse to use it for writing. You may use it for first drafts and refuse to use it for final versions. You may use it for technical questions and refuse to use it for questions that touch your judgment, your voice, or your conscience. The point of the framework is to give you the criteria for those decisions, not to make them for you.

The asymmetry worth knowing: AI is most useful for tasks where you already know what good output looks like and can correct it. AI is most dangerous for tasks where you don't, because you cannot tell when it has misled you. Use it for first drafts you will revise. Be cautious with it for final answers you will adopt.

Refusal as one option, not a moral high ground

If you choose not to use AI, that choice can be serious ethically. But the seriousness comes from understanding what you are refusing, not from the refusal itself. A refusal grounded in research is more rigorous than a refusal grounded in instinct, even when both feel equally principled.

The risk to watch for is treating non-use as proof of virtue rather than as a position you have argued for. Other students who use AI thoughtfully are not failing morally. Other students who refuse to use AI without understanding why are not succeeding morally. The work, in both cases, is in the reasoning.

A useful test: can you explain what specifically you are refusing? “I refuse to use AI” is a stance. “I refuse to use commercial generative AI for tasks that other methods can accomplish, because the environmental and labor costs are not justified by the educational gain” is a position. Positions are stronger than stances because they can be defended, refined, and shared.

If you refuse, write the AI Refusal Essay your professor may assign (it appears in the faculty section that follows). The essay asks you to explain what you refuse, why, what you still need to understand, and what policies you would support. That essay is itself a use of literacy, not a substitute for it.

3.2 For Faculty Teaching with AI

If you are a faculty member designing a course in which AI use is a question (and at this point that is most courses), the central design problem is not whether to permit AI but how to teach the discrimination that responsible use and responsible refusal both require. The strongest teaching principle in this framework is that students should not be forced into AI enthusiasm, but they should be invited into AI literacy. The five-layer framework in Part 2 (technical, environmental, ethical, practical, civic) gives you the curriculum vocabulary; the sections below give you assignments, statements, and rubric language.

If you are not convinced AI belongs in your course

Not every instructor reading this wants to teach with AI. Some are certain it should be kept out of the course entirely. That position deserves an argument rather than a workaround, and it has a serious advocate.

Weinreich (2026) makes the case for total opposition in research mathematics. His central claim is not that the output is poor. It is that a research paper has always certified something about the person who wrote it: that they practiced the discipline and arrived at understanding. Automating the production breaks the link between the practice and the thing that was measuring it. He asks which work a researcher understands best, and answers that it is the work they wrote themselves.

Move the claim into a classroom and it becomes a claim about assignments. An assignment certifies that a student practiced and arrived at understanding. If the student did not do the work, the assignment certifies nothing, whatever grade it earns.

This document agrees with that. Section 2.2 makes the same argument in the framework’s own terms: every literacy described here is built through practice that AI use can displace. The objection is not a misunderstanding to be corrected. It is correct.

The disagreement is about what follows from it.

One remedy is prohibition. Weinreich (2026) proposes, among several departmental measures aimed mainly at protecting research practice, that departments establish policies against AI in student work. In a course where students already use these tools daily, prohibition does not produce non-use. It produces undisclosed use. The student still offloads the effort, still receives the fluent output, and now also has a reason to conceal the whole transaction from the one person who could have shown them what it cost. The instructor gets the formation loss anyway and gives up the ability to teach against it.

Prohibition and protection look like the same instrument and are not. Prohibition forbids, and depends on detection to mean anything. Protection guarantees, and depends on nothing except how the assignment is built. An instructor who requires that a problem set be attempted unaided before any tool is opened has protected the formative work without needing to catch anyone. An instructor who bans the tool outright has protected nothing and now has a policing problem.

So the course this framework describes is closer to the skeptic's goal than a ban would be. It sets aside unaided work deliberately, says why, and then uses the tool in the open where its failures can be examined. The skeptical instructor was trying to preserve the conditions under which students learn. That is the same thing Section 2.2 identifies as the framework's precondition.

Two further things follow, and both favor the unconvinced instructor.

The first is that your skepticism is a teaching asset rather than an obstacle to work around. The practical literacy of a critic is sharper than that of an enthusiast, for the reason given in Part 2: the critic has spent more time looking at what is wrong with the output. An instructor who finds these tools unimpressive is unusually well equipped to show students exactly where they fail.

The second is that none of this requires you to recommend AI to anyone. The framework asks that students understand what these systems do, what they cost, and where they break. It does not ask you to endorse them. A course can teach AI literacy thoroughly while its instructor remains of the view that the technology should not have been built.

Designing for the five-layer framework

Each of the five literacies introduced in Section 2.1 maps onto a kind of teaching decision.

Technical literacy is built primarily through explanation and example. You do not need to teach the mathematics. You need to teach what a language model is doing well enough that students stop being surprised when it produces a confident wrong answer. A single lecture, paired with one assignment that requires students to find and analyze a specific AI error, will move most students from "AI is magic" to "AI is a probability machine with characteristic failure modes."

Environmental literacy is built through scale. Show students the IEA and OECD numbers (Part 1, Section 1.2). Then walk them through the Google single-prompt figure. The pedagogical aim is the moment when a student says, "So my prompts don't matter much, but the procurement decisions of my university do." That moment is environmental literacy.

Ethical literacy is built through cases. The standard ethics frameworks (UNESCO, OECD) are useful as scaffolding but insufficient as practice. Students develop ethical literacy by analyzing specific AI deployments in hiring, education, health care, policing, or creative industries, identifying which ethical concerns are most at stake in each.

Practical literacy is built through use and reflection. Students need experience with AI that wastes their time as well as AI that helps. The assignment templates below are designed for this.

Civic literacy is built through institutional connection. Have students read a real institutional document: a university AI policy, a state data-center regulation, an EU AI Act provision. Then have them identify what the document does and does not say, and what they would change. Civic literacy without engagement with real institutional text remains abstract.

Three participation modes

A course that requires students to use commercial AI tools forces a choice on students whose principled refusal has weight. A course that prohibits AI use leaves students unprepared for the institutional reality. The middle path is to offer multiple modes of participation that all develop the framework.

Active-use mode: Students use AI for selected tasks, document their prompts, evaluate the output, and explain what they accepted, rejected, or changed. This mode develops technical and practical literacy fastest.

Critical-observer mode: Students do not personally prompt AI but analyze instructor-provided outputs, evaluating hallucinations, bias, environmental cost, logical errors, or potential misuse. This mode preserves the option of principled refusal while building technical, ethical, and practical literacy.

A defense of this mode is warranted because it makes a deliberate trade-off. Practical literacy in this framework is the discrimination between AI uses where verification effort is worth the time and uses where it is not. That discrimination can be developed by observing a sufficient range of AI outputs and learning to predict where the failures will be; it does not require generating those outputs oneself. What Critical-observer mode does not develop is prompt-craft, which is a separate competency some students will need and others will not. Students whose participation mode is Critical-observer should be told this explicitly so the trade-off is theirs, not the instructor's.

Minimal-use mode: Students use AI only in controlled demonstrations, then write structured reflections on whether the use was justified. This mode emphasizes ethical and civic literacy.

Students choose their mode in consultation with you, with the understanding that the goal of the course is literacy, not compliance with a particular use pattern. A student who completes the course in critical-observer mode and can articulate a defensible position has met the course's objectives.

Assignments that teach AI literacy through skepticism

The following assignment templates were developed for environmentally concerned students but generalize to other principled-skepticism positions.

AI Footprint Audit: Students compare the environmental costs of different AI uses (text, image, video, model training, repeated prompting, large-scale deployment). The output is a one-page brief that identifies which use patterns are environmentally defensible and which are not.

Usefulness vs Cost Reflection: After a specific AI use (their own or an instructor-provided example), students evaluate whether the use was worth its environmental and educational cost. The point is not to reach a particular conclusion but to develop the habit of asking the question.

AI Ethics Case Study: Students analyze an AI application in a specific domain (hiring, education, health care, policing, climate modeling, art) and identify the three or four ethical concerns most at stake. The case study format prevents the dilution that comes with treating all ethical concerns as equally relevant in all cases.

Prompt Minimalism Exercise: Students solve a task with the fewest possible prompts, then compare the result against the same task done with careless iterative prompting. This exercise develops

practical literacy and environmental literacy together.

AI Refusal Essay: Students who choose not to use AI explain what specifically they refuse, why, what they still need to understand to defend the refusal, and what policies they would support. This is the assignment that signals to refusing students that their position is being taken seriously as a position, not as an exemption.

A useful design principle: assign AI as an object of study before assigning it as an assistant. The first AI assignment in your course should not be “use ChatGPT to draft something.” A stronger first assignment is: “Here are three AI outputs on the same question. Which is wrong? Which is wasteful? Which is most useful? What would you need to verify before trusting any of them?” That assignment changes the classroom dynamic. Students are asked to judge AI rather than to accept it.

Making environmental cost part of the rubric

When AI is used in your course, require students to account for it. A standard reflection question can be added to any assignment:

Was AI necessary for this task? Did it improve the work enough to justify its use? Could a smaller or less resource-intensive method have achieved the same result?

This reframes AI use as a decision with consequences rather than a default behavior. Three sentences in a rubric do more pedagogical work than a separate lecture on environmental ethics.

A model syllabus statement

The following statement can be adapted for the AI policy section of a syllabus:

Responsible AI Use and Environmental Concern. This course recognizes that artificial intelligence has real environmental, ethical, social, and educational costs. AI systems consume energy and water, depend on large-scale data infrastructures, raise concerns about privacy and authorship, and can be used in harmful ways. For that reason, this course will not treat AI as a neutral convenience or a substitute for learning. When AI is used, it will be used deliberately, sparingly, and critically. Students will ask whether the use is necessary, whether it improves learning, what risks it creates, and how its outputs should be verified. Students who object to using AI on ethical or environmental grounds may complete alternative forms of participation, such as analyzing instructor-provided AI outputs, writing critical reflections, or evaluating AI-related policy and environmental claims. The goal of the course is not to produce enthusiastic AI users. The goal is to produce informed, skeptical, ethically responsible thinkers in a world where AI is already widely deployed.

A model statement to say in class

In the first or second class meeting, a verbal version of the above signals the course’s stance directly:

I know some of you are concerned about AI's environmental impact and its ethical risks. I share many of those concerns. AI uses energy and water. It can be used to mislead people, invade privacy, weaken learning, imitate artists, and concentrate power. Those are real issues, and we will not brush them aside.

I want to separate two things: using AI carelessly and understanding AI carefully. You may decide that you do not want to use AI casually. That is a reasonable ethical position. But refusing to understand AI is different. AI is already being used in education, business, science, government, health care, and media. If thoughtful people decline to learn how it works, they will have less power to challenge it when it is used badly.

In this class, I will not ask you to trust AI. I will ask you to question it. I will not ask you to use it for everything. I will ask you to learn when it helps, when it harms, when it wastes resources, and when human judgment must take over.

3.3 For Administrators Setting AI Policy

If you are an administrator with responsibility for AI policy at a college or university (a provost, a dean, a CIO, a chief academic officer, a member of a faculty senate AI committee, a board member with technology portfolio responsibility), the question on your desk is not whether your institution will have an AI policy. You already have one, set by the cumulative effect of vendor contracts already signed, faculty practices already in place, and the absence of explicit positions on questions that are being asked anyway. The question is whether your institution will set its AI policy deliberately, or continue to let it be set by default.

What the framework asks of your institution

A curriculum organized around the five literacies in Part 2 requires several institutional commitments that are not always made explicit.

It requires that faculty have the discretion to make pedagogical judgments about AI use in their courses. That discretion is incompatible with mandates either to adopt or to prohibit AI universally. It is not incompatible with requiring AI literacy as an outcome. The difference between those two things is the subject of the section that follows.

It requires that students who refuse AI on principled grounds have viable participation paths in courses that engage AI as subject matter. That requires faculty support, alternative-assignment resources, and an institutional posture that treats principled refusal as a defensible position rather than as resistance to compliance.

It requires that the institution be able to answer the civic-literacy questions its students will learn to ask. What does the institution disclose about its AI vendor contracts? What environmental commitments has it made, and how do its AI procurement decisions sit with those commitments? What data terms govern the AI tools the institution provides? These questions are coming. The question is whether the institution will be prepared to answer them.

Requiring literacy without mandating use

Institutions increasingly want AI literacy as a stated learning outcome. The proposal draws principled objections from faculty and from students. Those objections deserve an answer rather than an exemption.

The answer begins with a distinction that is easy to state and often skipped. Requiring AI literacy is not the same as requiring AI use. They are different objects. One is an outcome: the student can explain what these systems do, what they cost, and where they fail. The other is a method: the student operates the tool. An institution can require the first without requiring the second, and the European Union's AI Act does exactly that. Article 4 places a literacy obligation on organizations that develop and deploy AI systems (European Union, 2024). A 2026 amendment softened it further, from ensuring a sufficient level of literacy to supporting its development, and stated plainly that no particular level is required of any individual. Through both versions the obligation runs to understanding. Neither obliges any particular person to use anything.

Section 2.2 makes the distinction more than a technicality. If the formation precondition holds, a literacy requirement that also mandated use would work against its own outcome. Stated properly, the requirement runs the other way. It obliges the institution to protect the unaided work.

That reframes both objections.

The faculty objection is that a requirement overrides professional judgment. It would, if the requirement reached the method. It does not. The requirement is on what students leave the course able to do. How they get there stays with the instructor, and Section 2.2 explains why it has to: only the person designing the assignments can sequence the unaided work that makes the outcome reachable. Discretion is not merely compatible with the requirement. The requirement depends on it.

The student objection is that a requirement compels engagement with a technology the student rejects. Here the distinction does most of the work. Understanding is not endorsement. A student may complete every element of AI literacy and conclude that the technology should not be used, and that conclusion is better supported for having been reached with knowledge. The strongest critics of any technology tend to be the people who understand how it works.

This is why the framework needs no waiver. A waiver exempts a student from an outcome. Nothing in the outcome is what the objecting student objected to. The objection is to use, and use is where discretion already lives.

There is a harder version of the objection, and none of the above answers it. Weinreich (2026) argues that individual choice in this area is shaped by conditions no individual controls. Once a practice becomes widespread, declining it stops being free and starts costing something. A framework that equips people to choose well may be answering a question the environment has already closed. He is right about the diagnosis. An institution that affirms discretion on paper has affirmed nothing if declining carries a professional penalty.

His remedy is for departments to act collectively against the technology: policies against its use in student work, hiring lines reserved for those who refuse it, more weight given in promotion to work produced without it. That is one way to change the incentive structure. It asks the institution to bet on an answer.

There is a narrower instrument. An institution can guarantee that declining to use these tools

carries no professional or academic penalty, and can say so in the documents where penalties are actually decided: assignment policies, evaluation criteria, promotion standards. That shifts the same incentive without requiring anyone to agree about whether the technology is good. It holds whether the skeptics turn out to be right or wrong, which is the property an institutional policy most needs while the question is open.

Policy questions to address

The following questions can structure an institutional AI policy review. Each has a default answer in the absence of an explicit policy, and the default answer is rarely the one the institution would choose if asked.

1. *Faculty discretion.* Does your institution affirm that individual faculty have the authority to set AI policy in their own courses, within broad institutional limits? Or does it leave that authority ambiguous?
2. *Student refusal.* Does your institution affirm that students may refuse AI use in coursework on principled grounds and complete alternative assignments? The Conference on College Composition and Communication has affirmed this right at the discipline-organization level (Conference on College Composition and Communication, 2026); is the institution prepared to back it?
3. *Vendor selection.* Who decides which AI tools the institution makes available, and on what criteria? What are the data terms? What is disclosed publicly about contracts? What environmental reporting is required from vendors?
4. *AI use in administrative functions.* The institution itself is an AI user, in admissions, grading support, advising, and student-success analytics. Is the policy on institutional AI use as stringent as the policy expected of faculty and students?
5. *Equity of access.* Premium AI tools cost money. If AI use becomes implicitly expected in coursework, students who cannot afford premium access are disadvantaged. What is the institution's position? (Extended treatment in the Equity, infrastructure, and procurement subsection below.)
6. *Disclosure.* When AI is used in evaluating student work (in detection, in grading support, in admissions screening), is that use disclosed to students? Should it be?
7. *Faculty development.* What support does the institution offer faculty who are designing courses around AI literacy, regardless of whether the courses require AI use?
8. *Regulatory readiness.* The EU AI Act's Article 4 literacy requirement applies to organizations that deploy AI systems. US state-level AI regulation is increasing. Has the institution mapped its AI deployment against current and forthcoming regulatory requirements?

Engaging with the professional society context

Professional societies are taking positions on AI in education at an accelerating pace. The Conference on College Composition and Communication passed a resolution in March 2026 affirming the right of students and faculty to refuse generative AI in the writing classroom (Conference on College

Composition and Communication, 2026). The American Historical Association's guiding principles articulate a position closer to engagement: that ignoring AI will neither halt its spread nor shield the discipline from its reach, while affirming that some educators have rejected AI for its ethical, environmental, and economic consequences (American Historical Association, 2025). EDUCAUSE's 2024 framework articulates a third position emphasizing literacy and competency development across institutional roles (EDUCAUSE, 2024).

These positions are not contradictory. They are responses to the same situation from organizations with different membership and different professional stakes. The administrator's question is whether the institution's policy is coherent with the positions its faculty's professional societies are taking. A writing department whose national organization has affirmed the right to refuse AI cannot reasonably be required to mandate AI use in writing courses, and an institutional policy that fails to recognize this puts the institution into avoidable conflict with its faculty.

The opposite case also holds. A history department whose national organization argues for engagement cannot reasonably refuse to develop AI literacy in its students, and a policy that does not support that engagement underserves the discipline.

The administrator's task is to know which positions the institution's professional societies have taken and to design institutional policy that respects those positions while preserving institutional coherence across departments.

Equity, infrastructure, and procurement

Three institutional questions sit beneath the visible policy questions.

Equity of AI access. Commercial AI tools have a free tier and a premium tier. If course assignments come to assume premium-tier capabilities, students whose families cannot afford the subscription are quietly disadvantaged. The institution has three real options: provide institutional access (which is expensive and creates vendor lock-in), prohibit assignments that require premium tiers (which constrains pedagogy), or accept the inequity as a cost of letting individual faculty decide (which the institution should be honest about if it is the chosen course). Avoiding the question does not produce a fourth option.

Data terms and institutional liability. When students use AI tools through institutional accounts, the institution becomes party to the data terms. What is disclosed in those terms about training data use, retention, and downstream sharing? When a faculty member uses an AI tool to evaluate student work, what data has been transmitted, to whom, and under what protection? These are not abstract questions. They are FERPA and contractual questions that will eventually be asked by counsel, by accreditors, or by litigation.

Environmental commitments and AI infrastructure. Many institutions have made formal sustainability commitments. AI deployment, particularly at scale, sits in tension with those commitments. The institution should be able to explain how its AI use is consistent with its sustainability statements, and where the inconsistencies are, what is being done about them. A student who has developed environmental literacy under the framework in Part 2 will eventually ask. The institutional answer should not be improvised.

A worked example: disclosing AI vendor contracts

A framework that asks institutions to be ready for civic-literacy questions should model what readiness looks like. The example below is a structure, not a template. It addresses one question, “What does the institution disclose about its AI vendor contracts?”, at the level of specificity an environmentally and ethically literate student would reasonably expect, and at the level of caution institutional counsel would reasonably impose.

A civic-literacy-ready disclosure on a public webpage would identify, at minimum: (1) the AI tools the institution has procured for instructional or administrative use, named by vendor and product; (2) the procurement decision-makers for each contract, by office or committee rather than by individual; (3) the data terms in summary form, specifically what student or employee data the vendor receives, retains, may use for model training, and may share downstream; (4) the environmental commitments the vendor has made and the institution’s basis for accepting or questioning those commitments; (5) the procurement cycle, so faculty and students know when contracts are up for review and what input is sought; and (6) what the institution has chosen not to disclose, with the reason (counsel-protected terms, vendor confidentiality clauses, pending negotiations) named explicitly.

The sixth item is the test. An institution that discloses items 1 through 5 but does not name what it is withholding has not modeled civic literacy. It has modeled selective transparency, which is the framing the document’s framework teaches students to recognize and reject. Naming the withheld categories does not require disclosing them. It requires acknowledging that the disclosure is not complete and saying why. This is a higher bar than most institutions currently meet, and it is also the bar the framework’s civic literacy layer prepares students to ask for. The mismatch is the work the framework asks the institution to do.

Communicating with students and faculty

A clear institutional communication about AI policy, written for both audiences, accomplishes more than a series of memos. It signals that the institution has thought about the questions, has taken positions, and has a process for revising those positions as conditions change.

The communication does not need to settle every question. It needs to be honest about which questions the institution has answered, which it is working on, and which it is leaving to faculty discretion. The middle category (questions actively being worked on) is the one most institutions try to avoid acknowledging. Acknowledging it builds trust. Pretending the questions have been answered when they have not erodes trust faster than any explicit policy position would.

A useful test: can the institutional AI communication be read by an environmentally concerned student without producing the conclusion that her concerns have been dismissed? If not, the communication needs revision.

4 Synthesis

The teacher’s role is not to persuade skeptical students that AI is good. The administrator’s role is not to procure AI tools because every institution is doing so. The student’s role is not to use AI in proportion to its perceived inevitability. The shared role of all three is different: to develop the

framework for asking better questions, about necessity, cost, accountability, design, and what should never be automated. The framework in Part 2 is not a set of answers. It is a way to organize the questions so that the answers can be defended, refined, and shared. Each audience holds different responsibility for those answers. The questions are common to all of them.

The questions that organize this document are not “Should we love AI or hate it?” They sort into three domains. On use, the questions are practical and immediate: When is AI worth using? When is it wasteful? When does it weaken learning? On accountability, the questions reach further: Who pays the environmental cost? Who profits, and who is harmed? What kind of transparency should institutions require, and how would noncompliance be sanctioned? On the civic-educational stakes, the questions are the ones that make the framework’s whole point: What should never be automated? What must humans still know how to do? What would a responsibly governed AI infrastructure look like, and who has the literacy to argue for it? These are the questions a course organized around the five literacies prepares students to ask. They are also the questions a faculty member’s syllabus and an administrator’s policy must engage.

The student in the opening scene is not a problem to be managed. She is the version of an AI-literate citizen the field most needs. If higher education engages her seriously, in courses that respect her refusal, with frameworks that develop the technical and civic literacy her skepticism deserves, and within institutions that have the policy courage to back her position, then her doubts become the field’s strongest critical asset. If higher education accommodates her into silence, or argues her into enthusiasm, the cost is not only to her education. It is to the public’s capacity to govern a technology that is already deciding how essential systems will be built. The best outcome of the framework in this document is not that every student uses AI. The best outcome is that the student who refused to use it becomes exactly the critic the field needs.

References

- American Historical Association. Guiding principles for Artificial Intelligence in history education, August 2025. URL <https://www.historians.org/resource/guiding-principles-for-artificial-intelligence-in-history-education/>. Approved by AHA Council, July 2025; published 5 August 2025.
- Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)*, pages 610–623. Association for Computing Machinery, 2021. doi: 10.1145/3442188.3445922.
- Cecilia Ka Yuk Chan and Wenjie Hu. Students’ voices on generative AI: Perceptions, benefits, and challenges in higher education. *International Journal of Educational Technology in Higher Education*, 20(1):43, 2023. doi: 10.1186/s41239-023-00411-8.
- Conference on College Composition and Communication. 2026 resolutions, March 2026. URL <https://cccc.ncte.org/cccc/2026-resolutions/>. Resolution affirming the right of students and teachers to refuse generative AI in the writing classroom. Annual Business Meeting, Cleveland, OH.
- Kate Crawford. *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. Yale University Press, New Haven, 2021. ISBN 978-0-300-20957-0.

- EDUCAUSE. AI literacy in teaching and learning: A durable framework for higher education, October 2024. URL <https://www.educause.edu/content/2024/ai-literacy-in-teaching-and-learning/executive-summary>. EDUCAUSE working group paper, 17 October 2024. Prepared by the AI Literacy Programs for Faculty, Staff, and Students Working Group; authored by Michelle Kassorla, Maya Georgieva and Allison Papini.
- European Union. Artificial intelligence act, article 4: AI literacy, 2024. URL <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>. CELEX: 32024R1689. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence.
- Yizhou Fan, Luzhen Tang, Huixiao Le, Kejie Shen, Shufang Tan, Yueying Zhao, Yuan Shen, Xinyu Li, and Dragan Gašević. Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance. *British Journal of Educational Technology*, 56(2):489–530, 2025. doi: 10.1111/bjet.13544.
- Google Cloud. Measuring the environmental impact of AI inference, August 2025. URL <https://cloud.google.com/blog/products/infrastructure/measuring-the-environmental-impact-of-ai-inference>. Blog post, 21 August 2025, accompanying the technical report on Gemini serving footprint.
- Veronika Hackl, Alexandra Elena Mueller, and Maximilian Sailer. The AI literacy heptagon: A structured approach to AI literacy in higher education. *Computers and Education: Artificial Intelligence*, 10:100540, 2026. doi: 10.1016/j.caeai.2026.100540. arXiv:2509.18900.
- Karen Hao. *Empire of AI: Dreams and Nightmares in Sam Altman’s OpenAI*. Penguin Press, New York, 2025. ISBN 978-0-593-65750-8.
- International Energy Agency. Energy and AI. Technical report, International Energy Agency, Paris, April 2025. URL <https://www.iea.org/reports/energy-and-ai>. World Energy Outlook Special Report.
- J. Paul Liu and Rachel Levy. Beyond tool adoption: A practical five-stage developmental continuum for AI literacy in higher education, 2026. URL <https://arxiv.org/abs/2606.00038>. arXiv:2606.00038. arXiv preprint. North Carolina State University.
- Duri Long and Brian Magerko. What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems (CHI ’20)*, pages 1–16. Association for Computing Machinery, 2020. doi: 10.1145/3313831.3376727.
- Davy Tsz Kit Ng, Jac Ka Lok Leung, Samuel Kai Wah Chu, and Maggie Shen Qiao. Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2:100041, 2021. doi: 10.1016/j.caeai.2021.100041.
- OECD. OECD AI principles, 2024. URL <https://oecd.ai/en/ai-principles>. Adopted May 2019; updated 3 May 2024, OECD Ministerial Council Meeting.
- OECD. The hidden costs of AI: Unpacking its energy and water footprint, February 2025. URL <https://oecd.ai/en/hidden-costs-ai>. OECD.AI Policy Observatory summary of the OECD/IEEE session at the French AI Action Summit, 12 February 2025, Paris.

OECD/European Union. Empowering learners for the age of AI: An AI literacy framework for primary and secondary education. Technical report, OECD Publishing, Paris, 2026. URL <https://ailiteracyframework.org/>.

Neil Selwyn. On the limits of artificial intelligence (AI) in education. *Nordisk Tidsskrift for Pedagogikk og Kritikk*, 10(1):3–14, 2024. doi: 10.23865/ntpk.v10.6062.

Emma Strubell, Ananya Ganesh, and Andrew McCallum. Energy and policy considerations for deep learning in NLP. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics (ACL 2019)*, pages 3645–3650, Florence, Italy, 2019. Association for Computational Linguistics. doi: 10.18653/v1/P19-1355.

UNESCO. Recommendation on the ethics of artificial intelligence, 2021. URL <https://unesdoc.unesco.org/ark:/48223/pf0000381137>. Adopted by acclamation by 193 Member States, 41st General Conference, 23 November 2021; published Paris: UNESCO, 2022, document SHS/BIO/PI/2021/1.

Max Weinreich. The crisis of AI-generated mathematics, August 2026. URL <https://arxiv.org/abs/2608.02859>. arXiv:2608.02859. arXiv preprint, math.HO, MSC 00A30.