

Teaching with Artificial Intelligence

MODULE 1

Foundations: Grounded Curiosity and Your Disciplinary Anchor

Faculty Development Reading

*Begin with what your discipline must protect, and most of the
technology questions begin to answer themselves.*

— the course stance, in one sentence

June 2026

Module Purpose

Most professional development about artificial intelligence starts in the wrong place: with the software. It demonstrates a tool, marvels at the output, and leaves you to figure out what any of it means for your Tuesday afternoon seminar. This course refuses that starting point. Almost nothing in teaching with AI is really about software; it is about learning, assessment, trust, and disciplinary judgment. Tool demonstrations age out in months. The judgment is what lasts (Module 1 lectures).

Module 1, therefore, begins with your discipline, not with a chatbot. By the end of this module you will draft the first two artifacts in your course portfolio: a **Discipline Statement** (250–350 words) naming what AI literacy means in your field, and a signed **Skeptic’s Charter** naming what AI must not damage in your teaching and what you would accept as evidence of damage or non-damage. The major decision you are preparing to make is a values decision, not a technology decision: *which intellectual practices in my discipline must remain substantially human-owned, and where might AI responsibly support access, practice, or feedback?* Every later module — the friction audit, the assessment decisions, the assignment redesign, the syllabus policy — will lean on the answer you give here.

Learning Objectives

After completing this reading and the Module 1 activities, you will be able to:

1. Articulate a personal stance toward AI in your teaching that is anchored to one disciplinary commitment AI puts under pressure and one that AI may sharpen.
2. Explain, in plain language a colleague could verify, what a large language model is doing when it generates text — and why fluency is not understanding.
3. Distinguish technical, applied, and critical AI literacy, and classify a course goal into one of those tiers.
4. Identify at least three access, language, disability, privacy, or workload conditions in your context that should constrain your later policy and assignment decisions.
5. Complete and sign a Skeptic’s Charter naming what AI must not damage in your teaching and what would count as evidence either way.

Before You Read: A 60-Second Retrieval Primer

This course opens every module by asking you to write before you read — a retrieval-practice move that makes the reading stick better than highlighting ever will. Take 60 seconds. No looking anything up.

YOUR 60-SECOND RESPONSE — write before you read

In 60 seconds: what do you currently believe AI is doing when it generates text? One or two sentences, in your own words. You will return to this answer at the end of the reading — and again at the end of the module.

Core Reading

A pedagogy course, not a tool tour

The stance of this course is **grounded curiosity**. Grounded, because every claim should be anchored in how learning actually works and in what your discipline actually values. Curious, because refusing to investigate a technology your students are already using is not a neutral act either. Grounded curiosity refuses the two framings that dominate most AI conversations in higher education.

The first is *hype*: AI as an oracle that will transform every syllabus, personalize every lecture, and liberate faculty from drudgery. The second is *panic*: AI as a cheating machine that has contaminated every assignment, to be met with bans and detection software. Both framings share a flaw — they make the technology the protagonist of your course. This course keeps your discipline, your students, and your judgment in that role. As the opening lecture puts it, the work is to make *defensible learning-design choices*: decisions you could explain to a skeptical colleague, a department chair, or a student who asks why the rules are what they are.

Grounded curiosity is also why this course asks you to build things. Across six modules you will produce six teaching-ready artifacts — a Discipline Statement and Skeptic’s Charter, a Friction Audit, an Assessment Stress Test with an AIAS Decision Memo, a Redesigned Assignment, a Syllabus AI Policy with an Equity and Ethics Audit, and a final portfolio with a Six-Month Implementation Plan. These are documents that go into a real syllabus in a real next term, not reflections about technology in general.

Skepticism is a stance, not a stage

If you are skeptical about AI in education, this course treats that as professional judgment doing its job — not as a developmental stage you are expected to outgrow. You may worry that students will outsource the work that builds judgment. You may worry that polished language will cover thin thinking, or that students with better tools, more time, or more confidence will gain an unearned advantage. You may simply be tired of being told that every new technology must now become part of your course. Those concerns are legitimate, and the research literature on academic integrity increasingly treats faculty mistrust as a stance to engage rather than overcome (Eaton, 2023).

What this module asks is one move: convert broad skepticism into specific, inspectable commitments. “AI is bad for learning” is hard to act on. “AI must not weaken my students’ ability to reason from primary sources, and I would accept declining performance on in-class source analysis as evidence of

weakening” is something you can design around, monitor, and defend. That conversion is exactly what the Skeptic’s Charter, described below, asks you to write down.

What a language model is actually doing

You do not need a technical background to make every pedagogical decision in this course. You need one functional idea: a large language model (LLM) is trained on enormous amounts of text and learns to predict which word — technically, which *token* — is most likely to come next, given everything that came before. That is the mechanism. There is no fact-checker inside, no database of verified claims, no understanding in the human sense. There is prediction.

One short example shows both the power and the problem. Ask a model to complete “The capital of France is —” and it will almost certainly produce *Paris*, because that continuation dominates its training data. But the same mechanism will fluently complete a question about a fictional country’s capital, a nonexistent journal article, or a fabricated court case — confident, well-formed, and citing nothing. The output is generated the same way whether it happens to be true or false.

Three practical consequences follow for teaching. First, **hallucination**: models produce fluent, confident output that is factually wrong, including invented citations. Second, **bias**: training data carries the biases of its sources, which is to say, much of the internet. Third, **overconfidence** — ours, not the machine’s: fluency reads as authority to most readers, including strong students. None of these is a bug to be patched out of your course planning; each is a design condition to plan around.

Two more ideas complete the functional picture. Models display a **jagged frontier**: they handle tasks unevenly, writing a passable sonnet while failing to count the letters in a word, and there is no clean line that predicts which side of the frontier a task falls on (Mollick, 2024). And all of it is subject to **capability drift**: what models can and cannot do has changed materially every six months for several years. Any specific claim about AI capability in this course — including this paragraph — should be read with its vintage stamp: *capability claims in this reading are valid as of June 2026 and should be re-verified afterward*.

That is the entire model explanation, and the course stops here on purpose. You will not be taught prompt tricks or tool interfaces. You will be taught to choose postures and design boundaries — work that survives the next model release.

Fluency is not understanding

The single most useful sentence in this module for your grading practices is this: **fluency is not understanding**. For a language model, fluent text is the direct product of prediction, not a signal of comprehension. The problem is that for the whole history of your discipline, fluent text from a student was a reasonable proxy for thinking — imperfect, but workable. Generative AI quietly breaks that proxy. A beautifully organized explanation of a concept now tells you that fluent text about the concept exists in the world, not that the student who submitted it can explain, apply, or extend the concept. The course returns to this in Module 2 as the Safety Gap. For now, hold the implication: when AI is in the picture, evidence of learning has to be designed, not assumed.

PAUSE AND APPLY — one task, held in mind

Pick one task from one of your courses — a close reading, a lab report, a case analysis, a problem set, a critique, a clinical reflection. Hold it in mind for the rest of this reading; it becomes your test case in the activity.

Now ask: *if a student completed this task with substantial AI assistance, what would I still know about their learning?* Write your one-sentence answer in the margin. Do not solve the problem yet — just notice the size of it.

Oracle, Adversary, Tutor: the posture question

When faculty imagine AI in a course, two default postures dominate. The **Oracle** posture trusts the tool to be right: ask it anything, accept the answer, let it settle the question. It is risky precisely because judgment is what many of our courses exist to teach, and a fluent answer is not a defensible answer. The **Adversary** posture inverts the error: it treats every student as a suspect and designs the course from suspicion — detection software, interrogations, shrinking trust. It feels protective, but it narrows teaching toward surveillance and punishes the innocent along with the guilty.

This course refuses both defaults and centers a third posture: **AI as Tutor** — not an always-correct expert, but a bounded support for practice, questioning, rehearsal, feedback, or revision. Sometimes useful, often limited, never a substitute for the core intellectual work. The practical question of the whole course is not “how do I defeat AI?” or “how do I adopt AI?” but “*which posture, with which boundaries, serves this learning goal — if any?*” Module 2 develops these postures as design choices; Module 4 turns them into assignment language. The table below is your first reference version.

Posture	What it assumes	Primary risk in a course	How this course responds
Oracle	The tool’s answer can be trusted and judgment can be delegated to it.	Students skip the schema-building work; fluent answers pass for defensible ones.	Treat AI output as a claim to be verified, never as a settled answer.
Adversary	Students will cheat by default; the course must be designed from suspicion.	Surveillance displaces teaching; trust and help-seeking collapse; false accusations fall unevenly.	Design for evidence of learning instead of detection of misconduct.
Tutor	AI can support practice, questioning, and feedback within explicit boundaries.	Without boundaries, the Tutor quietly becomes an Oracle that does the work.	Choose the posture deliberately, bound it explicitly, and keep the core work student-owned.

Table 1. Three postures toward AI. The course’s stance — grounded curiosity — refuses the first two as defaults and treats the third as a bounded design choice.

Three tiers of AI literacy — and why they are disciplinary

“AI literacy” risks becoming the new “critical thinking”: invoked everywhere, defined nowhere. This course uses a working three-tier model. **Technical AI literacy** is understanding enough about how these systems work to recognize their limits — prediction, training data, failure modes. **Applied AI literacy** is using AI for a real task with awareness of those limits. **Critical AI literacy** is judging where AI should not be used, where it causes harm, and who governs it. The international AILit Framework behind this model describes AI literacy across four domains — knowledge of AI, durable skills with AI, ethical judgment about AI, and orientation toward AI’s social effects (European Commission & OECD, 2025).

The course’s distinctive claim is that AI literacy is a *disciplinary* responsibility, not a generic add-on. Each field has to define what understanding, using, questioning, and governing AI means inside its own intellectual tradition. In the humanities, AI puts pressure on voice, evidence, judgment, and dialogue. In the sciences, on replicability, method, and falsifiability. In professional fields, on ethics of practice and responsibility to clients, patients, and communities. Your course almost certainly already implicates applied or critical AI literacy without naming it — a nursing course teaching calibrated clinical judgment is teaching exactly the capacity an Oracle posture erodes.

COMMON TRAP — the “critical thinking” deflection

When asked what AI must not damage, the most common faculty answer is “critical thinking.” As written, it is a deflection: too broad to design around, too vague to monitor, and identical across every discipline — which means it names nothing about *yours*.

The fix is to replace the slogan with the specific cognitive practice. Not “critical thinking” but “discriminating between two historical sources whose framing differs,” “choosing a statistical method and defending the choice,” “translating a client’s vague request into testable requirements.” Your Discipline Statement will be evaluated on exactly this specificity. (A related trap: “AI is just another calculator.” The analogy captures routinization, but the calculator does not generate plausible falsehoods — or essays.)

Equity before policy

Most AI guidance treats equity as a closing paragraph. This course surfaces it in Module 1 because access conditions should *constrain* your design decisions, not decorate them. Five conditions matter most. **Access:** AI tools differ sharply between free and paid tiers, and students without funds, devices, or reliable internet experience a different technology than their classmates. **Language:** AI can be a genuine support for multilingual students — translation, rephrasing, reading-level adjustment — and the same students are disproportionately flagged by AI-detection tools. **Disability:** AI offers real accessibility affordances and also fails unevenly for users of assistive technology. **Privacy:** students should never be required to feed personal data into commercial systems as a condition of coursework. **Workload:** any redesign you adopt has to be feasible for the actual staffing of your course — including contingent and adjunct colleagues teaching large enrollments without course releases (Furze, 2026).

EQUITY AND ACCESS CHECK

When you write your Discipline Statement, name the student you are presuming — and the student you are not. Picture one absent or marginal student in your course: the working parent on a phone, the multilingual writer, the student using a screen reader, the student who cannot afford a subscription.

Then carry one equity concern forward in writing. The worksheet asks for it explicitly, and your Module 5 policy audit will check whether the concern survived into your course rules. An AI stance that only works for fully resourced students is not a stance; it is a filter.

The Skeptic's Charter and the Discipline Statement

Module 1 closes with two short writing tasks that anchor the entire portfolio. The **Skeptic's Charter** is a one-page document you write and sign. It answers three questions: What must AI not damage in my teaching? What would I accept as evidence of damage? What would I accept as evidence of non-damage? The third question is what makes the Charter intellectually honest — it commits you to noticing if your fears do *not* materialize, just as the second commits the optimists among us to noticing if they do. You will not edit the Charter during the course. You will reopen it in Module 6 and write a response: what was honored, what was challenged, what you still hold.

The **Discipline Statement** is a 250–350 word position statement built from the activity's scaffolding: three student capabilities in one course (one knowledge, one process, one judgment capability), each classified by AI-literacy tier; one capability that must remain substantially human-owned; one place AI might responsibly support access, practice, or feedback; and one equity concern you will carry forward. If you do not have a current course, the course provides five Synthetic Course Packets — a humanities seminar, an introductory statistics course, a nursing ethics course, a composition course, and a studio art course — and you may adopt one as your working context for all six modules.

Connection to the Module Artifact**CONNECTION TO YOUR PORTFOLIO ARTIFACT — Discipline Statement + Skeptic's Charter**

This reading gave you the four raw materials the Module 1 worksheet asks for: a functional understanding of what LLMs do (so your statement is grounded, not mystical); the posture vocabulary (so you can say what role AI may and may not play); the three literacy tiers (so you can classify your course goals); and the equity conditions (so your stance works for your actual students).

Open the *Discipline Values and AI Literacy Snapshot* worksheet and the *Skeptic's Charter* worksheet in the Module 1 folder. Draft the statement first; let it inform the Charter. Then check your draft against the list below — it mirrors the rubric the course steward will use.

Before you submit, confirm:

- ☐ I named at least one *specific cognitive practice* — no unmodified “critical thinking” anywhere in the statement.
- ☐ I listed three capabilities (knowledge, process, judgment) and classified each as technical, applied, critical, or not primarily AI literacy.
- ☐ I identified one capability that must remain substantially human-owned, with a reason a colleague could challenge.
- ☐ I identified one place AI might responsibly support access, practice, or feedback — with a boundary.
- ☐ I named one equity concern I will carry into Modules 3–5.
- ☐ My Skeptic’s Charter answers all three questions — including what I would accept as evidence of *non-damage* — and is signed and dated.

Key Terms

Term	What it means in this course
Grounded curiosity	The course stance: refuse both AI hype and AI panic; investigate AI as a learning-design problem anchored in disciplinary values.
Large language model (LLM)	A system trained on vast text that generates output by predicting the most likely next token given everything before it. Prediction, not understanding.
Hallucination	Fluent, confident output that is factually wrong or fabricated — including citations to sources that do not exist.
Fluency is not understanding	The course’s shorthand for why polished AI-era text no longer functions as a proxy for student thinking.
Jagged frontier	AI’s uneven capability profile: strong on some tasks, surprisingly weak on adjacent ones, with no clean predictive line (Mollick, 2024).
Capability drift	The pace at which AI capabilities change — faster than syllabi normally revise. The reason every capability claim carries a vintage stamp.
Technical AI literacy	Understanding enough about how AI systems work — prediction, training data, failure modes — to recognize their limits.
Applied AI literacy	Using AI for a real task with awareness of its limits and appropriate documentation.
Critical AI literacy	Judging where AI should not be used, where it causes harm, and who governs it.
Oracle / Adversary / Tutor	Three postures toward AI: trusted answer machine; presumed cheating threat; bounded support for practice and feedback. The course refuses the first two as defaults.

Term	What it means in this course
Skeptic's Charter	A signed one-page document naming what AI must not damage in your teaching and what counts as evidence of damage or non-damage; reopened in Module 6.
Discipline Statement	A 250–350 word position statement defining what AI literacy means in your field and what students must be able to do that AI cannot responsibly do for them.

Brief Applied Example

A history professor drafts her anchor. Professor Reyes teaches a sophomore survey in U.S. history with 32 students, about a third of them multilingual. Her first instinct on the worksheet is to write that AI must not damage “students’ critical thinking.” The trap box stops her, and she rewrites: the capability her course actually builds is *weighing two primary sources whose accounts of the same event conflict, and explaining which to trust for which question*. That is a judgment capability, and she classifies it as critical AI literacy: a model can summarize both documents fluently, but the act of adjudicating between them — and owning the argument — is the discipline.

Her statement then names a knowledge capability (placing events in causal sequence — technical tier, low AI risk) and a process capability (locating and citing archival sources — applied tier, where AI may help with search strategies but invents citations often enough that verification must stay student-owned). For responsible support, she notes that several students read nineteenth-century prose at frustration level; AI rephrasing of *secondary* material at an accessible reading level, followed by work with the original, is a support she is willing to design for. Her equity concern: those same multilingual students must never be exposed to AI-detection accusations on the basis of their prose style. Her Charter records what she would accept as evidence of damage — declining in-class source analysis over two terms — and of non-damage: stable or improving analysis even as take-home polish rises. She signs it, dates it, and moves on, mildly surprised that the module never asked her to open a chatbot.

End-of-Reading Reflection

1. Return to your 60-second primer response. How would you revise your description of what AI is doing when it generates text? Which single word in your original answer now looks most misleading?
2. For your own course — or your chosen Synthetic Course Packet — name one disciplinary commitment AI puts under pressure and one it might sharpen. What evidence would tell you, a year from now, which way it went?
3. The course argues that skepticism is a stance to engage, not a stage to outgrow. Where does your skepticism (or your enthusiasm) most need evidence rather than instinct? What would change your mind?

Sources and Further Reading

The module materials cited throughout this reading — the Module 1 lectures, the Module 1 Participant Workbook and Glossary, the Discipline Statement samples library, and the Synthetic Course Packets — are provided inside the course.

Boussioux, L. (n.d.). *AI with Leo: A hands-on guide to deep learning and large language models*.
<https://aiwithleo.lovable.app>.

Eaton, S. E. (2023). Post-plagiarism: Transdisciplinary ethics and integrity in the age of artificial intelligence and neurotechnology. *International Journal for Educational Integrity*, 19.

European Commission, & OECD. (2025). *Empowering learners for the age of AI: An AI literacy framework for primary and secondary education* (Review draft, May 2025). With Code.org.

Furze, L. (2026). *Teaching AI ethics: A guide for educators*. Leon Furze. <https://leonfurze.com>.

Harvard University. (n.d.). *AI teaching resources*. <https://www.harvard.edu/ai/teaching-resources/>

Mollick, E. (2024). *Co-intelligence: Living and working with AI*. Portfolio/Penguin. Chapter 1.

Stanford Teaching Commons. (n.d.). *Artificial intelligence teaching guide*. Stanford University.
<https://teachingcommons.stanford.edu/teaching-guides/artificial-intelligence-teaching-guide>

Suggested optional extension: Mollick’s Chapter 7 (“AI as a Tutor”) previews the learning-science questions Module 2 takes up; the Stanford guide’s self-paced AI literacy modules pair well with this module’s three-tier model.