

Teaching with Artificial Intelligence

MODULE 5

Syllabus AI Policy, Equity, Ethics, and Course Norms

Faculty Development Reading

A policy is what a tired student understands at midnight before the deadline. Everything else is decoration.

— the course stance, in one sentence

June 2026

Module Purpose

You now hold a redesigned assignment with an explicit AI contract. Module 5 scales that clarity to the course level: a syllabus policy that tells every student, for every assignment, what AI use is acceptable, what is not, and what to do when they are unsure or have made a mistake. The module begins by clearing the ground: it rejects detection-led integrity practice — not because integrity stops mattering, but because detection is unreliable, inequitable, and corrosive to the trust a course runs on. In its place it builds policy from transparent expectations, student responsibility, evidence of learning, and equity-aware implementation (Module 5 lectures).

The portfolio artifacts are a **Syllabus AI Policy** paragraph (150–250 words) built from a seven-component template, an **Equity and Ethics Audit** of that policy across eight dimensions, one first-week asynchronous touchpoint, and one revised sentence — highlighted, because the audit is only real if it changes something. The major decision you are preparing to make: *what are my course-level AI rules, and do they survive contact with my least-resourced student?*

Learning Objectives

After completing this reading and the Module 5 activities, you will be able to:

1. Explain in writing why AI detection is an unreliable and inequitable foundation for academic integrity practice, and name at least three replacement evidence sources.
2. Define the core intellectual substance of one of your courses — the thinking, judgment, and final claims students must own.
3. Draft a syllabus AI policy paragraph using the seven-component template, anchored in that substance and including a help-seeking invitation.
4. Conduct an equity and ethics audit of the policy across eight dimensions — access, language, disability, privacy/data, truth/verification, copyright, labor/environment, and power/bias — and revise at least one sentence in response.
5. Plan one first-week asynchronous touchpoint that builds AI norms without requiring live attendance.

Before You Read: A 60-Second Retrieval Primer

YOUR 60-SECOND RESPONSE — write before you read

Write the current AI-use sentence from your syllabus — or write “none.” Then read it aloud as if you were a tired student. What would that student assume they can and cannot do? Sixty seconds. Most current policies fail this little test in instructive ways; keep your answer for the reflection at the end.

Core Reading

Why detection fails as a teaching strategy

Start by separating four claims that get blurred together in integrity conversations (Module 5 lectures). *AI-content detection* — software scoring text as machine-generated — is the unreliable one. *Plagiarism detection* matches text against sources; it is somewhat reliable, but answers a different question. *Authorship evidence* (stylistic analysis) is suggestive at best. *Process evidence* — what your course design itself produces — is the one this course builds on. The case against leading with the first claim has three layers, in ascending order of importance.

The reliability problem. Detection accuracy degrades under paraphrasing and basic prompt variation, vendor claims have not matched independent audits, and capability drift means today's detectable output is tomorrow's undetectable output — in both directions (claims dated June 2026; see also Furze, 2026, on truth and post-plagiarism). A score is not a finding. **The inequity problem.** Detection tools disproportionately flag multilingual writers, neurodivergent writers, and students who use AI legitimately for accessibility, such as translation and sentence restructuring. The students most exposed to false accusations are precisely the students with the least institutional power to contest them. **The trust problem — the deepest.** When students learn that a course relies on detection, they stop asking questions about AI use. Help-seeking is suppressed; the integrity culture collapses into cat-and-mouse. A course cannot simultaneously invite students to discuss their AI use and threaten them with software that scores it. This course's position is careful: detection scores are never *standalone* evidence for integrity action, which is different from pretending misconduct doesn't exist. The replacement is a stack of better evidence: process evidence built into assignments (Module 4); low-stakes in-course conversation about AI use; comparison of submitted work against demonstrated capability (your Cognitive Sanctuary, quietly doing double duty); and self-disclosure built into the documentation norm.

When expectations are violated anyway — and sometimes they will be — the course follows Eaton's restorative line: the first response is developmental, a structured conversation about what happened and what the expectation was, with sanctions available but not leading (Eaton, 2023). The policy-language version of all this replaces "any AI-generated work is plagiarism" with something a student can actually use: *submitting work as your own when you did not do the cognitive work it was meant to demonstrate is academic dishonesty — and each assignment names what that cognitive work is.*

Core intellectual substance: the anchor of the policy

A policy needs something to protect, or it is just a list of prohibitions. The course names that something **core intellectual substance (CIS)**: the thinking, judgment, and final claims for which the student is responsible. CIS is disciplinary — in a writing course, it is the argument; in a quantitative course, the interpretation; in a clinical course, the calibrated judgment. You have been building this answer since Module 1: your Discipline Statement named it, your Friction Audit located it in one assignment, your AIAS memo ruled on it. The policy now states it at course scale, which is why the strongest policy

paragraphs read less like terms of service and more like a sentence from your Discipline Statement with rules attached.

Three ways policies fail

Reading flawed policies is the fastest way to learn the genre — the module’s markup exercise hands you a deliberately broken one. The failure modes to spot: **the blanket ban with detection threat** (fragile for every reason above, and it teaches students the course’s primary stance toward them is suspicion); **anything-goes permission with no documentation** (reads as generous, collapses the integrity floor, and quietly advantages students with paid tools and prior AI fluency); and **vague “responsible use” language with no examples** (transfers all interpretive risk to the student — and the students who guess wrong will not be a random sample). Vagueness is not neutrality; it is a tax on the students with the least insider knowledge.

The seven components of a defensible policy

The course’s policy template assembles seven short components, each doing one job (Module 5 Policy Anatomy card). Together they fit in 150–250 words — a paragraph students will actually read, not a contract they will scroll past:

1. **Stance** — one sentence on your course’s relationship to AI, anchored in what the course teaches.
2. **Rationale** — one sentence connecting the stance to learning, not to enforcement.
3. **General expectation** — what is permitted and prohibited at course level, in concrete verbs.
4. **Documentation expectation** — what students record when they use AI; two sentences maximum.
5. **Privacy warning** — what should never be entered into AI tools: personal data (theirs or others’), confidential material, unpublished work they do not own.
6. **Help-seeking invitation** — where to ask when unsure, and what happens if they make a mistake; the sentence that keeps questions above ground.
7. **Assignment-specific note** — the pointer that AIAS levels vary by assignment and each assignment’s AI Use box governs.

Tone is a design property of the genre, not a flourish: the policy is written for an actual student who may be tired, anxious, or unsure — a thoughtful colleague writing to a thoughtful student, not a lawyer writing to a defendant. The module’s **read-aloud audience test** operationalizes it: read your draft aloud as if you were the student, and answer three questions. Does this invite questions? Does it raise or lower the floor of help-seeking? Does it tell the student what to do if they have already made a mistake? A policy that fails the third question manufactures concealment at precisely the moment honesty would have been cheap.

COMMON TRAP — the “responsible use” mirage

“Students are expected to use AI responsibly and ethically” feels finished and governs nothing. Responsible by whose definition? For which task? Documented how? The sentence reads as policy to the instructor and as a riddle to the student — and it will be interpreted most generously by the students with the most confidence, and most anxiously by the students with the least.

The repair is always the same: replace the adverb with a scenario. *“You may use AI to brainstorm and to check your grammar; you may not use it to draft your analysis; when you use it, add one sentence saying how.”* If a sentence in your draft cannot be translated into a scenario, it is decoration — cut it or convert it (Module 5 Policy Anatomy card, before/after pairs).

The Equity and Ethics Audit: eight dimensions

A drafted policy is a hypothesis. The audit tests it across eight dimensions — and ethics here is wider than the usual pair of bias and paid access. The framing follows Furze’s *Teaching AI Ethics* (2026), which treats truth, copyright, privacy, data, environment, labor, and power as teachable, auditable dimensions rather than disclaimers. For each row, the worksheet asks: OK, Needs Revision, or Not Applicable — and at least one Needs Revision must end in an actual revised sentence.

Dimension	The audit question you ask of your policy
Access	Can every required AI use be completed on a free tier, on a phone, on a slow connection? Is there a documented alternative for students who cannot or choose not to use AI?
Language	Can a B2-level English reader parse every sentence? May students use AI to translate or rephrase course materials — and does the policy say so?
Disability	Is the policy screen-reader friendly? Does it leave room for AI as assistive technology without forcing disclosure of disability status?
Privacy and data	Does the policy warn students what never to paste into AI tools? Does any assignment effectively require surrendering personal or third-party data?
Truth and verification	Where AI-produced information is allowed, does the policy require verification against named sources — with citation?
Copyright	Does the policy address attribution when AI output is included, and students’ responsibility for material AI was trained on or reproduces?
Labor and environment	Are AI’s human-labor and environmental costs named where relevant — not as a mandated position, but as acknowledged context?
Power and bias	Whose perspectives does the system amplify or mute? Where AI touches content about marginalized groups, is a bias acknowledgment required?

Table 1. The eight-dimension Equity and Ethics Audit (Module 5 worksheet; dimensions grounded in Furze, *Teaching AI Ethics*, 2026). Each dimension is marked OK / Needs Revision / Not Applicable — and at least one revision must be made and highlighted.

Two dimensions deserve a sentence more than the table can give. *Truth and verification* is where your discipline's epistemics enter the policy: an AI summary is a claim about sources, not a source, and the policy should say who is responsible when a confident fabrication slips through (the student who cited it — which is why verification must be taught, not assumed). *Labor and environment* is the dimension faculty most often skip as political; the course's requirement is deliberately modest — name the costs where relevant, require no position — because pretending the supply chain doesn't exist is itself a position (Furze, 2026).

EQUITY AND ACCESS CHECK

The audit's sharpest test is cumulative: read your policy one last time as the multilingual student working 30 hours *and* on a phone's free tier. Each dimension you marked OK in isolation may still stack against this student. The course floor — no required paid tools, no detection scores as standalone evidence, a stated alternative path, a help-seeking sentence — exists because policy failures compound on exactly the students with the least slack to absorb them.

Then check the mirror image: does the policy *permit* the support these students legitimately need? A policy silent on translation, rephrasing, and assistive use will be read as prohibiting them, by precisely the students too cautious to ask.

First-week norms: the policy is a conversation, not a clause

A policy read once in week one and never touched again will not survive week six. The module, therefore, ends with one first-week asynchronous touchpoint — a low-stakes structure that makes students handle the policy rather than scroll through it. The template pack offers four formats: a **scenario quiz** (five short student-facing situations, each answered with an AIAS level and one sentence of rationale); an **anonymous question form** (with a promised synthesis back to the class — the instructor's answers become the policy's living FAQ); a **policy quiz** (five low-stakes items with explanations); or a **reflective scenario** (one 150-word situation, three open prompts). All four run without live attendance, all four work in any LMS, and all four lower the cost of asking — which, after everything this module has argued, is the actual mechanism of integrity (Module 5 First-Week Templates; cf. Harvard AI Teaching Resources on co-created norms).

Alignment: the policy, the scale, and the assignment must agree

Last check before drafting: your three layers of AI communication now have to tell one story. The syllabus policy explains the framework and the floor; the AIAS vocabulary names levels; each assignment's AI Use box rules its components. A policy that says "AI is generally discouraged" above an assignment that integrates an AI Adversary critique is not nuance; it is a contradiction students will resolve in whatever direction maximizes survival. Run the three layers together once — policy paragraph, quick-reference scale, one redesigned assignment — and read for daylight between them. Where the layers disagree, the assignment is usually right, because it was designed against an outcome; revise the policy toward it.

PAUSE AND APPLY — draft the two hardest sentences

Before opening the worksheet, draft components 1 and 6 — the stance sentence and the help-seeking invitation. The stance: one sentence beginning “In this course, AI...” that a student could repeat accurately to a roommate. The invitation: one sentence telling a student who is unsure — or who has already crossed a line — exactly what to do next, without flinching.

Read both aloud in your tired-student voice. If the stance sounds like a vendor or a prosecutor, or the invitation sounds like a trap, revise before the template makes the rest easy.

Connection to the Module Artifact

CONNECTION TO YOUR PORTFOLIO ARTIFACT — Syllabus AI Policy + Equity/Ethics Audit + first-week touchpoint

The Module 5 worksheet sequences the work exactly as this reading did: paste your current policy (or “none”); draft the seven components with character budgets and examples; add the assignment-specific note for your Module 4 redesign with its AIAS level; run the eight-dimension audit; revise and highlight at least one sentence; choose and build one first-week touchpoint from the template pack.

The markup exercise — a deliberately flawed sample policy failing at least three audit dimensions — is the calibration step; do it before auditing your own draft, because flaws are easier to see in someone else’s syllabus. The samples library and the six-question critical-friend protocol (multilingual sticking points, hidden burdens, threat-reading, paid-tool assumptions, documentation bloat, missing privacy warnings) complete the self-review.

Before you submit, confirm:

- ☐ All seven components are present and the whole policy is 150–250 words — short enough to be read, specific enough to be used.
- ☐ The stance and rationale connect to my course’s core intellectual substance, not to enforcement or fashion.
- ☐ Permitted and prohibited uses are scenario-concrete; no unconverted “responsible use” sentences remain.
- ☐ Documentation expectation is two sentences or fewer; privacy warning names what never goes into the tools.
- ☐ The help-seeking invitation tells an unsure (or already-erring) student exactly what to do next.
- ☐ The eight-dimension audit is complete; at least one sentence is revised and highlighted, with its 30-word justification.
- ☐ Policy, AIAS levels, and my redesigned assignment tell one consistent story; my first-week touchpoint is chosen and drafted.

Key Terms

Term	What it means in this course
Detection-led policy	Integrity practice built on AI-content detection scores. Rejected here as unreliable, inequitable, and corrosive to help-seeking — never standalone evidence.
Replacement evidence stack	Process evidence, in-course conversation, comparison with demonstrated capability, and self-disclosure — the course's substitute for detection.
Restorative integrity	Responding to apparent violations developmentally first — conversation before sanction — with sanctions available but not leading (Eaton, 2023).
Core intellectual substance (CIS)	The thinking, judgment, and final claims the student must own in your course; the anchor every policy component points back to.
Seven-component policy	Stance, rationale, general expectation, documentation, privacy warning, help-seeking invitation, assignment-specific note — in 150–250 words.
Help-seeking floor	The policy property that determines whether unsure students ask or hide; raised by invitation, lowered by threat.
Free-tier floor	The course rule that every required AI use must be completable on a free tier, on a phone, on a slow connection.
Read-aloud audience test	Reading the policy aloud as a tired student: Does it invite questions? Raise help-seeking? Say what to do after a mistake?
Equity and Ethics Audit	The eight-dimension review — access, language, disability, privacy/data, truth/verification, copyright, labor/environment, power/bias — ending in one real revision.
B2-readable	Parsable by an upper-intermediate English reader; the course's clarity floor for all student-facing policy language.
First-week touchpoint	A low-stakes asynchronous structure (scenario quiz, anonymous question form, policy quiz, or reflective scenario) that makes students handle the policy.
Datafication	The conversion of student work and behavior into data for systems students did not choose — the reason the privacy warning and audit rows exist (Furze, 2026).

Brief Applied Example

A business instructor audits her own draft. Professor Imani teaches a marketing capstone in which teams build campaign strategies for real local clients. Her draft policy is friendly and AIAS-aligned: Level 3 for most deliverables with a decision record, Level 1 for the individual strategy defense. Then the audit starts flagging rows. **Privacy and data:** her assignments require analyzing client briefs — real businesses' unreleased plans — and nothing in the policy stops a student from pasting a client's brief into a free AI

tool that may retain it. Needs Revision. **Access:** her “use the AI tool of your choice” sentence quietly assumes paid tiers; two of her teams discovered last term that the free tier’s usage caps broke their workflow mid-sprint. Needs Revision. **Copyright:** AI-generated imagery in student campaigns was attributed nowhere. Needs Revision — one sentence now requires labeling AI-generated assets and recording the tool.

Her highlighted revision is the privacy sentence: “*Client materials are confidential: never enter a client’s documents, data, or identifying details into any AI tool; use the anonymized case template instead.*”

Thirty words of justification, and a workflow change she would never have caught from the integrity angle alone. Her first-week touchpoint is the scenario quiz, rewritten with marketing cases — including one scenario where the *right* answer is to ask her first. The read-aloud test reshapes her final sentence from “violations will be pursued” to “if you are unsure, or you realize you have crossed a line, come talk to me — that conversation is always the better path.” Her tired student, she decides, can work with that.

End-of-Reading Reflection

1. Return to your primer sentence (or your “none”). Rewrite it as a stance-plus-invitation pair. What did your original sentence teach a tired student that you did not intend to teach?
2. Which audit dimension surprised you most when you turned it on your own course or Synthetic Course Packet — and what is the one sentence you will revise in response?
3. The module claims help-seeking is the actual mechanism of integrity. What in your current course raises the cost of a student asking an AI question out loud — and what would lower it in week one?

Sources and Further Reading

The module materials cited throughout this reading — the Module 5 lectures, the Module 5 Participant Workbook, the Policy Anatomy reference card, the Sample Policy for Markup, the First-Week Touchpoint template pack, and the policy samples library — are provided inside the course.

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

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/>

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

Suggested optional extension: Furze’s chapters on power and datafication extend audit rows seven and eight beyond the syllabus into curriculum design; the Stanford policy workshop materials include additional sample policies worth marking up with the Module 5 protocol. Detection-reliability claims in this reading are stated as of June 2026 and should be re-verified before reuse.