

Teaching Through the Machine

*Protect the thinking, reclaim your time,
and teach AI honestly.*

EDUCATORS EDITION

Dr. Ayse Ozturk

Clinical Associate Professor of Marketing · Darla Moore School of Business, University of South
Carolina

drayseozturk.org

We can't detect our way out of this. So let's teach.

I teach, and I sat where you're sitting: the first wave of AI-written work landing in the inbox, the urge to find a tool that would catch it, the slow realization that the tool would fail and take a few honest students down with it. The detector arms race is one we lose, and the students who pay for our false positives are too often the ones writing carefully in a second language.

So this edition is not about catching students. It is about the part of our job AI cannot touch — the thinking — and how to protect it, surface it, and assess it when every learner has a fluent draft a click away. It is written peer to peer, from one instructor's notes to another's. It spends its pages on five things: why detection is a trap and design is the answer, what AI still gets wrong about the work you assign, how to reclaim hours of grading and prep without lowering your standards, how to teach students to direct these tools instead of outsourcing their minds to them, and why AI literacy now belongs in your course no matter what you teach.

Everything here comes from my running notes at drayseozturk.org, checked against primary research, with sources linked so you can take any claim to a colleague — or a committee. Where the evidence is early, I say so. Read it in one sitting, then change one assignment before the next term.

Ayse

Dr. Ayse Ozturk · Columbia, South Carolina

Five moves for the AI classroom

Each chapter ends with three things you can change before the next class.

01 The detection trap — design over policing.

AI detectors are unreliable and biased against your most careful students. The durable fix isn't a better detector; it's an assignment worth doing.

02 What AI still gets wrong about your syllabus.

It can pass the exam and still invent a citation. Knowing where its competence ends is what makes it safe to use in your course prep.

03 Grade faster without grading worse.

AI can draft the rubric, the first-pass feedback, the quiz bank. The judgment it can't supply — knowing this student — is the teaching.

04 Teach students to direct AI, not outsource thinking.

Offloading the work feels like learning and isn't. The skill worth grading is using the tool in a way that still builds the mind.

05 AI literacy is the new critical reading.

Bias, fabrication, cost, limits — what every student should be able to interrogate by graduation, in any discipline you teach.

01

CHAPTER ONE

The detection trap — design over policing.

Every hour spent trying to prove a student used AI is an hour not spent making that question worth answering for themselves.

There are real fingerprints in AI prose — the balanced chiasmus, the asyndetic tricolon, the tidy rule of three, the over-fondness for the em-dash, a register that stays formal when a person would relax. You can learn to see them. The trouble is twofold: the fingerprints fade with every model, and the machines built to spot them are wrong often enough to be dangerous.

A Stanford study put numbers on it. Leading detectors classified essays by native English speakers almost perfectly, then flagged more than 60% of TOEFL essays — written by non-native speakers — as AI-generated. The reason is mechanical: careful writers who use plainer, more standard syntax produce exactly the low-surprise text these tools read as artificial. The detector doesn't catch cheating. It catches difference.

A detector flagged 61% of essays by non-native English writers as fake. It wasn't finding AI. It was finding accents.

OpenAI quietly retired its own AI-text classifier for low accuracy. That is the company that makes the model telling you the detector doesn't work. Build your academic-integrity policy on that, not on a vendor's confidence interval.

61.3%

Average false-positive rate when leading GPT detectors scored TOEFL essays by non-native English writers as "AI-generated." Near-zero for native US essays.

[Stanford HAI, 2023](#)

The tells, for reference.

Chiasmus · the rule of three · asyndetic lists · em-dash chains · relentlessly formal tone. Useful signals — never proof.

Move the question, not the goalposts

If detection can't be the wall, the assignment has to be. The most robust response is not surveillance; it is designing work that a generic AI draft can't satisfy — because it asks for something the model doesn't have. Anchor the task in this week's in-class discussion, a local dataset, the student's own earlier draft, a primary source you chose. Ask for the reasoning, the dead ends, the revision history, not just the polished output. Make process visible and assessable, and a frictionless final draft stops being the prize.

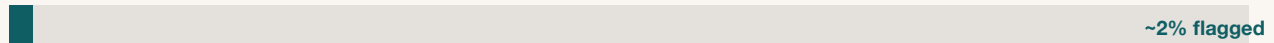
This is not extra work so much as different work, and it has a side benefit: assignments that resist AI tend to be the more authentic, more durable assignments anyway. The question worth asking is not "did they use AI" but "was this worth doing without it."

Reframe the whole

question. "Did a student use AI?" is unanswerable and divisive. "Does this task reward thinking the AI can't do?" is in your control.

Detector verdict on human-written essays

Native English writers



Non-native English writers (TOEFL)



Same tool. Same human authorship. The only variable is who wrote more "standard" prose.

FIGURE 1. The cost of a false positive isn't evenly distributed. Detectors disproportionately accuse the careful, the second-language, the formula-following student — the people least able to absorb the accusation. Source: Stanford HAI, 2023.

PUT IT TO WORK

- 1 Retire the detector from anything consequential.** If a tool can't tell a TOEFL essay from a bot, it cannot be the basis for an integrity charge. Use judgment and conversation, not a score.
- 2 Rewrite one major assignment to require something AI lacks** — in-class material, a personal data source, an annotated process, an oral defense of the work.
- 3 State your AI policy as expectations, not threats.** Tell students what use is welcome, what must be disclosed, and why — clarity prevents more misconduct than fear does.

CHAPTER TWO

02

What AI still gets wrong about your syllabus.

It can pass your final and then fabricate a source for the answer.
Knowing where its competence stops is what makes it safe to lean on.

Al capability is jagged, not smooth. Robotician Hans Moravec noticed the pattern long before chatbots: the high-reasoning tasks we treat as hard are comparatively easy to automate, while ordinary common sense is not. A model can ace a bar exam and then fail to reliably alphabetize a list. For a teacher, that means competence on one task tells you almost nothing about the next one.

Two failure modes matter most in a classroom. First, confident fabrication: models produce the most plausible next words, and plausible is not true. They invent citations, misquote sources, and assert dates with total composure. Tellingly, the newer "reasoning" models don't fix this — on OpenAI's own factual benchmark, its o3 and o4-mini hallucinated more than the older model, not less. Second, currency: a model's built-in knowledge has a cutoff and a slant, so it can be quietly out of date or skewed on exactly the contested topics your course exists to examine.

It can pass the final and **invent the citation.** Trust it task by task, never wholesale.

The lesson isn't distrust. It's calibrated trust: use AI where the work is easy for it and cheap for you to check, and keep your hand on the wheel everywhere else.

**33–
48%**

How often OpenAI's newer o3 and o4-mini models hallucinated on a factual-accuracy benchmark — higher than the older model, not lower. "Smarter" did not mean more truthful.

[OpenAI system card, 2025](#)

Moravec's paradox.

Human-hard and machine-hard are different lists. Plan around what the model finds easy, not what you do.

A trust map for course work

Sort the tasks you might hand a model by one question: *can I verify the output in seconds?* Rephrasing a paragraph, generating practice problems in a form you can check, drafting discussion prompts, varying the wording of a quiz — all easy to verify, safe to lean on. Anything where the cost of a confident error is high and the check is slow — supplying facts and statistics, attributing quotations, summarizing a source you haven't read, judging the nuance of a borderline grade — stays under your hand or out of the model entirely.

The same rule you'll teach students applies to you first: never publish or grade on an AI claim you haven't independently confirmed. A fabricated citation is embarrassing in a student paper and worse in your slides.

The one test. "Can I check this in seconds?" If yes, lean on AI. If no, the model is a suggestion, never a source.

Verifiability of the output →



FIGURE 2. A trust map for the tasks on your desk. The dividing line is not difficulty but verifiability — how quickly you can catch a confident error before it reaches a student.

PUT IT TO WORK

- 1 Sort your prep tasks into the two columns** above before next term. Automate the left freely; keep the right under human review.
- 2 Verify every AI-supplied fact, quote, and citation** before it reaches a slide, a handout, or a grade. Treat the model as a fast intern, not a reference.
- 3 Show students a real hallucination from your field.** Nothing teaches calibrated trust faster than watching the tool confidently invent a source they can check.

CHAPTER THREE

03

Grade faster without grading worse.

The model can draft the rubric, the feedback, and the quiz bank. What it can't do — know this student, in this course — is the part that was always the teaching.

Here is the reframe that gives you your evenings back. AI gets you to roughly 70% of most teaching artifacts almost instantly — a rubric draft, a first pass of feedback on a paper, a bank of quiz items, a reading guide, a rewrite of a confusing prompt. The remaining 30% is the part that requires you: knowing that this student finally grasped causation and should hear it, that this cohort is lost on a concept the draft glides over, that this borderline case needs mercy or rigor. That 30% is not the leftover. It is the teaching.

Used this way, AI doesn't lower your standards; it moves your hours from production to judgment. You stop typing the same rubric for the fifth time and spend that time on the comments only you can write.

***AI can draft the feedback.
Knowing *which student needs
to hear what is the part that
was always the job.****

The shift in how you ask matters too. Don't script the model step by step. Define the outcome and the constraints — the level, the tone, the rubric, the misconceptions to target — and let it draft. The field calls this moving from prompting to **context engineering**: the win is feeding it your materials, not finding magic words.

70 / 30

AI drafts ~70% of a teaching artifact instantly. The 30% you add — knowing the learner, the cohort, the call — is irreducibly yours.

Context over cleverness.

Give the model your rubric, your reading, your standards. A grounded draft beats a clever prompt every time.

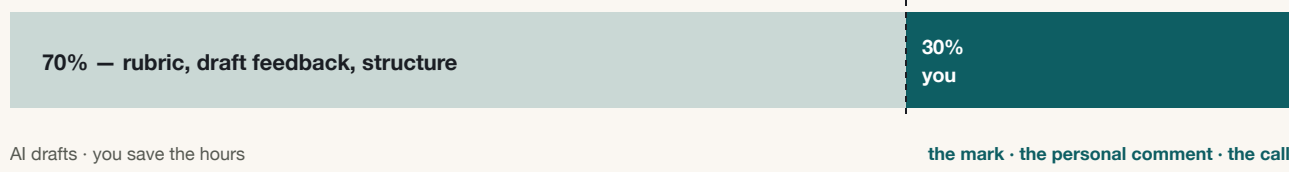
Keep the human in the loop you grade on

One firm line: AI can help you draft feedback, but the grade and the personal comment are yours. Let the model surface patterns across a stack of essays, or propose feedback you then edit, accept, or discard. Never let it assign the mark. Students can tell generic comments from real ones, and the trust you spend there is hard to earn back. The goal is to spend less time generating and more time deciding.

A fair word on confidentiality: don't paste identifiable student work into consumer tools whose data terms you haven't checked. Use institution-approved platforms, or strip names. Time saved is not worth a privacy breach.

Privacy line. Student work is protected data. Use approved tools or anonymize before any AI sees it.

A graded paper, returned



Spend the time you reclaim on the right-hand block — that's where the learning lands.

FIGURE 3. The 70/30 split, for grading. Automation collapses the cost of the left block; your value — and the student's trust — lives entirely in the right. Reclaimed time should flow there, not out the door.

PUT IT TO WORK

- 1 Draft your next rubric with AI, then edit.** Give it the assignment and your standards; refine the draft instead of writing from scratch. Bank the time.
- 2 Use AI for first-pass feedback, never the grade.** Let it propose comments you approve or rewrite; keep the mark and the personal note in your hands.
- 3 Set one privacy rule for yourself** — approved tools or anonymized text only — and follow it before any student work touches a model.

04

Teach students to direct AI, not outsource thinking.

Handing the work to a model feels productive and can leave nothing behind. The skill worth building is using the tool in a way that still grows the mind.

The danger isn't that students use AI. It's that using it can feel exactly like learning while leaving no trace behind. A 2025 MIT Media Lab study had people write essays with ChatGPT, with a search engine, or unaided, and tracked their brains with EEG. The AI group showed the weakest neural connectivity of the three — and most striking for any teacher, a large majority couldn't quote a single line from the essay they had just produced. The work got done. The learning didn't stick.

Treat that study with care: it was small, early, and not yet peer-reviewed, and its authors warned against the "AI makes you dumb" headlines it spawned. But the mechanism it points at is one every educator already knows. Effort is where learning happens. A tool that removes the effort can quietly remove the learning, and the student feels more capable precisely as they become less so.

Outsourcing the thinking feels like learning. Naming that gap for students is half the job.

So the goal is not abstinence; these tools are part of their future work. The goal is teaching the difference between AI that does the thinking for you and AI that makes you think harder — and grading for the second.

83%

of participants who wrote with ChatGPT couldn't quote a sentence from the essay they had just written. The unaided group could.

MIT Media Lab, 2025 (preprint, n=54)

Read it honestly. A small, early study — suggestive, not settled. Cite it as a prompt to think, not as proof.

Make the thinking the thing you grade

The fix is built into the assignment. Ask students to critique an AI draft rather than submit one — find its errors, its weak evidence, its missing counterargument. Have them show the prompts they used and what they changed and why. Grade the revision, the judgment, the argument they can defend out loud — the parts that require a working mind. Used as a sparring partner that students must out-think, AI raises the floor of the conversation. Used as a ghost-writer, it hollows the assignment out.

Name the trade openly with your students. Most aren't trying to cheat; they're trying to save time, and no one has told them that this particular shortcut spends something they'll want later. Teach them when reaching for the tool makes them sharper and when it just makes them faster at learning nothing.

The reframe for students.

Don't ask the AI to do the assignment. Ask it to give you something to argue with.

Cognitive engagement while writing (MIT Media Lab, EEG)

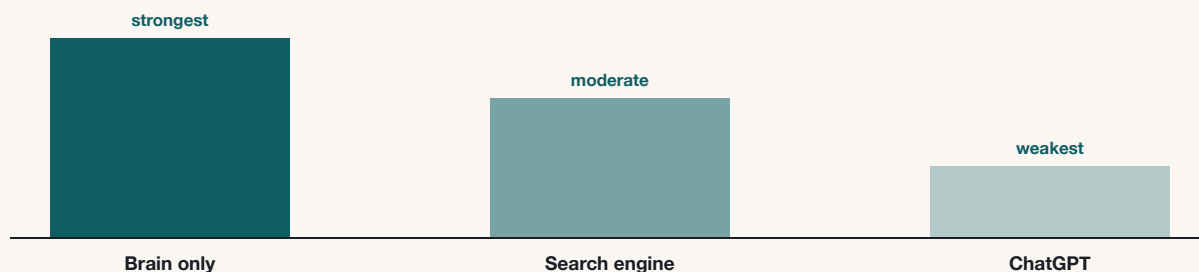


FIGURE 4. Relative brain connectivity across the three writing conditions, as reported by the MIT study (illustrative, not to scale). The more the tool did, the less the mind engaged. Early evidence — but a useful picture to keep in front of students. Source: MIT Media Lab, 2025.

PUT IT TO WORK

- 1 Assign critique, not generation.** Have students find the flaws in an AI draft and improve it — the errors are the lesson, and the work can't be outsourced.
- 2 Require a process trail.** Prompts used, what they kept, what they rejected and why. Grade the judgment, not just the product.
- 3 Have one honest conversation about cognitive offloading.** Show them the study. Let them decide where the shortcut costs them — they'll self-regulate better than any ban.

05

AI literacy is the new critical reading.

Whatever your discipline, students will graduate into a world that hands them confident machine-written answers. Teaching them to interrogate those answers is now part of your job.

For a generation we taught critical reading: who wrote this, what do they want, what's missing, can I trust it. AI literacy is the same skill aimed at a new kind of source — one that is fluent, tireless, confident, and wrong in ways that don't announce themselves. It is not a computer-science topic. It belongs in the history seminar and the nursing lab as much as the coding course, because every graduate will be handed machine-written answers and asked to act on them.

Four habits cover most of it. **Interrogate the output:** assume it may be fabricated until checked, and ask for sources you can follow. **See the bias:** these models carry the slants of their training and their makers — even vendors' own tests show measurable differences in how evenhandedly models treat contested topics. **Understand the limits and the costs:** a knowledge cutoff, a jagged competence, and a real if shrinking energy and water footprint behind every answer. **Disclose and attribute:** treat undisclosed AI use the way we treat any unattributed source.

AI literacy isn't a tech skill. It's critical reading for a source that never stops talking.

Bias is measurable. Even makers' own evaluations show models differ in how evenhandedly they treat political topics — a fact students should know, not assume away.

[Anthropic, 2025](#)

It's cross-curricular. The history, nursing, and business student each need this as much as the CS major. Don't outsource it to one department.

You don't need a new course

This rarely requires a new syllabus line — it requires threading the four habits through what you already teach. Bring a real AI answer into class and have students fact-check it against the reading. Compare how two models handle a contested question in your field. Add one line to your assignment sheet on how AI may be used and how it must be disclosed. Model your own checking out loud, so students see calibrated trust performed by an expert rather than preached. The aim is a graduate who reaches for these tools fluently and is never fooled by them.

That is the through-line of this whole edition. We don't protect learning by walling AI out of the classroom; the world won't. We protect it by teaching students to think with the machine and past it.

Smallest viable step. One fact-checking exercise and one disclosure line on the syllabus already moves a student toward literacy.

1 · Interrogate

Assume it may be wrong. Ask for sources. Check before you trust.

2 · See the bias

It carries its training's slant. Compare models on hard topics.

3 · Know the limits

Knowledge cutoff, jagged skill, and a real footprint per answer.

4 · Disclose

Attribute AI use as you would any other source.

FIGURE 5. The four habits of AI literacy — a portable framework that drops into any discipline without a new course. Interrogate, see the bias, know the limits, disclose.

PUT IT TO WORK

- 1 Run one fact-check exercise this term.** Bring a real AI answer to class; have students verify it against the reading and find what it got wrong.
- 2 Add a disclosure line to your syllabus.** State plainly how AI may be used in your course and how its use must be acknowledged.
- 3 Check something out loud.** Model your own verification of an AI claim in front of students — expert habits taught by demonstration stick.

Where to take this next

Four pages on drayseozturk.org, chosen for the questions teachers actually ask — each with why it earns your time.

AI for Educators

drayseozturk.org/ai-notes

The full running notes this edition draws from — the teaching-focused observations, kept current, when you want the source behind a claim or the latest thinking on a topic here.

Custom AI Teaching Tools

drayseozturk.org/resources

Ready-to-use tools and templates for course work — the practical companion to Chapter 3 when you're ready to draft a rubric or a quiz bank rather than read about it.

Prompting & Context

drayseozturk.org/resources-prompting

How to write the outcome briefs and context packs from Chapter 3, so the model drafts from your standards and materials instead of generic filler.

AI Dictionary

drayseozturk.org/ai-dictionary

Plain-language definitions for the terms your students (and committees) will use — tokens, hallucination, context window, agents — so you can teach them precisely.

SOURCES CITED IN THIS EDITION

Stanford HAI / Liang et al., [GPT detectors are biased against non-native English writers](#) (2023) · MIT Media Lab, [Your Brain on ChatGPT](#) (2025, preprint, n=54) · OpenAI, [o3 and o4-mini System Card](#) (2025) · Anthropic, [Measuring political bias in Claude](#) (2025). Model-specific figures are point-in-time, as measured on the cited benchmarks, and will move with each model generation; the MIT study is early and not yet peer-reviewed — cite it as suggestive.

DR. OZTURK EDITIONS

Short, rigorous field guides to working with AI.

Each edition takes the same body of notes and rebuilds it for one audience — educators, professionals, students, beginners, everyday life, and advanced practitioners — so you read only what changes your practice. Clear, concrete, current, and cited. Made to be finished in one sitting and shared with a colleague the next day.

drayseozturk.org

A FREE AI-LITERACY HUB BY DR. AYSE OZTURK · DARLA MOORE SCHOOL OF BUSINESS