

Reading the Detector: What False Positives Teach Us About AI Literacy

Weekly Analysis — <https://ainews.social>

A student submits an essay she wrote herself, at her own kitchen table, over three caffeinated nights. A piece of software scans it and announces, with the cool confidence of a parking meter, that it is sixty-eight percent likely to be machine-generated. The number is wrong. But the number is also, in a sense, the most honest teacher in the room — not because it tells the truth, but because of what it reveals about how little we understand the machine doing the telling. The false positive is a small disaster for the student and a large gift for the rest of us, if we are willing to read it correctly.

We are not, mostly, reading it correctly. Colleges have spent millions of dollars licensing AI-detection tools whose accuracy claims collapse under scrutiny, and they have done so while the tools' own makers quietly hedge their guarantees [9]. Teachers reach for these tools anyway, often knowing they are unreliable, because the alternative — sitting with uncertainty — is harder than outsourcing judgment to a dashboard [3]. This is not a story about cheating. It is a story about literacy: about what it means to live among systems that generate confident-sounding outputs, and whether we have built the civic and cognitive equipment to interpret them.

[9] Colleges pay millions for AI detectors that are flawed - CalMatters

[3] AI detection tools are unreliable. Teachers are using them anyway : NPR

The argument of this essay is that the false positive is not a defect to be engineered away. It is a curriculum. Every wrongful flag is a compressed lesson in probabilistic reasoning, in the limits of pattern-matching, in the difference between a guess and a verdict. The trouble is that we have organized our entire public conversation about "AI literacy" to skip that lesson — to treat literacy as a set of skills for *using* the tools rather than a capacity for *reading* them skeptically. What follows is an attempt to map the muddle: to ask what we mean by literacy, whose definition is winning, and what that winning definition leaves us unable to see.

The Two Literacies Fighting for the Word

"AI literacy" sounds like a single competence, the way "reading" sounds like one thing. It is not. At least two distinct projects are

wearing the same coat, and they want very different things from you.

The first project is operational. It treats literacy as fluency: the ability to prompt, to integrate, to deploy. Microsoft’s own widely circulated training path frames literacy as understanding “what AI is, how it works, and how to use it responsibly,” a definition that slides almost immediately from comprehension into capability [14]. Its companion course for educators is even more explicit about the destination — not critical distance but confident adoption, a workforce ready to put the tools to work [4]. In this frame, the literate citizen is a competent operator. Whole curricula now elevate prompt engineering to the status of a foundational “21st century skill,” as though the central question of our moment were how to phrase a request to a chatbot rather than whether to trust its answer [11].

There is nothing wrong with operational fluency as far as it goes. The difficulty is that it tends to colonize the word, leaving no room for the second project, which is interpretive. This second literacy treats the central skill not as *use* but as *judgment*: the capacity to ask what a system is actually claiming, on what evidence, with what margin of error, and in whose interest. UNESCO, in the materials it has built for educators, keeps insisting on the second kind — pointing out that generative systems are “stochastic in generating outputs,” producing different answers to identical inputs, and are therefore “potentially less trustworthy, especially for the teaching of factual and conceptual knowledge” [18]. That single sentence carries more genuine literacy than a semester of prompt-craft, because it tells you something about the *nature* of the thing rather than how to operate it.

Here is the move to watch. When an institution says it is teaching AI literacy, it is almost always teaching the first kind while borrowing the moral prestige of the second. The detector scandal is what happens when a population fluent in *using* tools has never been taught to *read* them. A teacher who can navigate a detection platform’s interface but cannot interrogate its false-positive rate is operationally literate and interpretively illiterate — and that combination is precisely the one that produces a wrongful accusation delivered with a straight face.

What the Detector Cannot Do, and Why That Is a Lesson

The deepest thing a false positive teaches is also the most mathematical, and the most resistant to slogans: no detector can be simultaneously highly sensitive and highly specific. This is not a flaw in any particular product. It is a structural tradeoff built into the act of clas-

[14] Introduction to AI Literacy - Training | Microsoft Learn

[4] AI for educators - Training | Microsoft Learn

[11] Frontiers | Prompt engineering as a new 21st century skill

[18] AI competency framework for teachers - UNESCO

sification itself. Tune the system to catch every machine-generated essay (high sensitivity) and you will inevitably sweep up human writing that happens to be clean, formulaic, or non-native (low specificity). Loosen it to spare the innocent and you let the actual machine text walk through. There is no setting that escapes the dilemma; there is only the choice of which error you would rather make.

The empirical record bears this out with embarrassing consistency. The Stanford-affiliated AI Index, which catalogs the field’s own published research, includes empirical studies of detector performance demonstrating how easily these tools are fooled and how unevenly they perform across populations [18]. Reporting on classroom deployment found teachers adopting detection tools at scale precisely as the evidence of their unreliability accumulated — a striking case of practice racing ahead of warrant [16]. And the harm is not randomly distributed. Analyses of detection in schools document that these tools disproportionately flag the writing of students who learned English as a second language, whose prose statistically resembles the smoothed-out patterns the detectors associate with machines [2]. The false positive, in other words, has a demographic shape. It falls hardest on the people least equipped to contest it.

To understand why this is inevitable rather than fixable, it helps to borrow a concept from machine learning’s own vocabulary. Janelle Shane’s accessible account of how these systems fail describes overfitting through the image of an AI that, having only ever seen roads bordered by grass, panics when it crosses a bridge — it learned a correlation and mistook it for a law [18]. A detector trained to associate certain rhythms of sentence and certain distributions of vocabulary with “machine” has overfit in exactly this way. It has not learned to recognize machine writing; it has learned to recognize a *texture*, and that texture is shared by plenty of human beings — the careful, the coached, the translated, the neurodivergent, the merely conventional. Shane’s larger point applies directly: everyday AI, she reminds us, “is not flawless, not by a long shot” [18]. The literate response to a detection score is therefore not “is it right or wrong” but “what texture is it actually responding to, and who shares that texture innocently?”

This is what I mean by the false positive as curriculum. It is a free, vivid, repeatable demonstration that classification is probabilistic, that probability has costs, and that those costs are paid by real people. A society that absorbed that lesson would have taken a serious step toward genuine literacy. Instead we have largely learned the opposite lesson — that the machine is a kind of oracle whose pronouncements describe reality rather than estimate it.

[18] HAI_AI-Index-Report-2024

[16] More Teachers Are Using AI-Detection Tools. Here’s Why That Might Be a ...

[2] AI Cheating in Schools: 2026 Global Trends & Bias Risks

[18] You Look Like a Thing and I Love You

[18] You Look Like a Thing and I Love You

The Tyranny of the Two-Way Sort

There is a second falsehood the detector teaches, subtler than the first and in some ways more corrosive. By design, it answers a binary question: human or AI. It hands you a number that resolves, in the user's mind, into one of two boxes. And that binary is a lie about how text now gets made.

Almost nothing in contemporary writing is purely one or the other. A student drafts a paragraph, a tool suggests a revision, she rejects two suggestions and accepts a third, then rewrites the whole thing in her own cadence. A journalist uses a model to summarize a transcript, verifies the summary against the source, and folds three sentences of it into an article that is otherwise entirely hers. The reality is a collaborative continuum, a spectrum of human-machine entanglement with countless gradations — and the detector's binary erases every one of them. It forces a smear of grey into black or white and then reports the result as though the grey had never existed. The literacy we actually need is the literacy to see the continuum the tool is built to hide.

This matters far beyond the schoolhouse, which is why it belongs at the center of any civic account of literacy. Consider how the same binary thinking distorts our politics. The fear of election deepfakes has produced a public primed to ask of any video, "real or fake?" — and that primed binary turns out to be exploitable in both directions. Researchers documenting the 2024 US election cycle found that the *actual* damage of synthetic media came less from undetected fakes than from the collapse of trust the fear of fakes produced [10]. The Brookings analysis of the same period warned, pointedly, that false *claims* of deepfakery — politicians dismissing genuine recordings as AI fabrications — were poised to do more harm than deepfakes themselves [21]. The careful study of the so-called Slovak case, in which a faked audio recording surfaced days before a vote, complicates the easy panic still further, suggesting the recording's actual influence was far murkier than the headlines implied — neither the clean catastrophe nor the clean hoax the binary demands [8].

The thread connecting the wrongly flagged essay and the disputed campaign recording is the same impoverished literacy: a public trained to want a verdict and a technology happy to supply one. When the only question you know how to ask is "human or machine, real or fake," you become equally vulnerable to the false positive that condemns the innocent and the false negative that launders the guilty. You also become trivially manipulable by anyone who grasps that "it's AI" has become an all-purpose solvent for inconvenient evidence. The

[10] Deepfakes in the 2024 US Presidential Election

[21] Watch out for false claims of deepfakes, and actual ... - Brookings

[8] Beyond the deepfake hype: AI, democracy, and "the Slovak case"

continuum-literate citizen asks a different and harder set of questions: How was this made? Through what chain of hands and tools? What would count as verification, and has anyone done it?

Reading Scores the Way We Read Polls

If the operational definition of literacy gives us competent operators, and the binary gives us bad bookkeeping, what would the interpretive literacy actually look like in practice? The most useful analogy is one the public has, imperfectly, already half-learned: the polling margin of error.

Most educated readers now know, at least in principle, not to treat a poll showing a candidate at forty-eight percent as a statement that the candidate *has* forty-eight percent. They know there is a margin, that the margin matters most when the race is close, that the methodology shapes the meaning, and that a single poll is a data point rather than a prophecy. This is hard-won statistical literacy, and it is exactly the disposition a detection score requires. A "ninety-two percent likely AI" reading is not a measurement of reality; it is an estimate with an error band, a confidence the tool projects rather than a fact it discovers. To read it literately is to ask about its margin, its base rate, its methodology, and its track record on people like the person being judged.

This is precisely where the operational framing fails the citizen, because it never teaches the statistical disposition at all. It teaches you to read the number as an output to act on, not an estimate to interrogate. The reporting on classroom misinformation makes the stakes vivid: when researchers examined what makes students — "and the rest of us," the headline pointedly adds — fall for AI-generated falsehoods, the vulnerability was not technical ignorance but a failure of probabilistic and source-critical habits, the same habits that would let you read a poll or a detection score with appropriate caution [22]. A systematic scoping review of generative AI and misinformation reached a convergent conclusion: the protective factor is not detection capability but the broader apparatus of critical evaluation — provenance, corroboration, source criticism — that long predates the technology and extends far beyond it [12].

UNESCO's work on artificial intelligence and disinformation lands in the same place, insisting that the answer to machine-generated falsehood is not a better machine but a more demanding citizen — one equipped with the verification reflexes of source criticism rather than the false comfort of an automated arbiter [13]. There is a temp-

[22] What Makes Students (and the Rest of Us) Fall for AI Misinformation?

[12] GenAI and misinformation in education: a systematic scoping ... - Springer

[13] Inteligencia artificial y desinformación - UNESCO

tation, well-funded and constant, to believe we can automate our way out of this — that the solution to bad AI is good AI. Some of that promise is real; tools that surface provenance and flag known fabrications genuinely help [6]. But notice what the detector scandal warns against: the same confidence we are urged to place in the fake-news fighter is the confidence that, misplaced, produces the wrongful flag. An automated arbiter you cannot interrogate is a liability whether it is accusing a student or adjudicating the news. The margin-of-error disposition is the one that lets you accept the help without surrendering the judgment.

[6] AI shows promise in the fight against fake news

The Demand the Literate Citizen Can Make

Probabilistic thinking and source criticism are dispositions a person can cultivate alone. The third skill the false positive teaches cannot be practiced in private, because it is a demand made on power: the demand for transparency. And here the conceptual map gets genuinely contested, because transparency is exactly what the operational definition of literacy never asks for.

A fluent operator does not need to know a detector's false-positive rate to use it; the interface works the same either way. But a literate citizen cannot do without it, and that information is routinely withheld. The institutions that spent millions on these systems frequently could not, when pressed, produce independent validation of the accuracy rates they were paying for [9]. The vendors treat their thresholds and training data as proprietary. The user is handed a number and denied everything that would make the number interpretable. To be literate, in the fullest sense, is to recognize that withholding as itself a power move — and to insist on the disclosure as a precondition of trust.

[9] Colleges pay millions for AI detectors that are flawed - CalMatters

Kate Crawford gives this move its sharpest name. In her account of the political economy of these systems, she identifies "enchanted determinism" — the discourse that urges us "to focus on the innovative nature of the method rather than on what is primary: the purpose of the thing itself," and that "obscures power and closes off informed public discussion, critical scrutiny, or outright rejection" [18]. The detection score is enchanted determinism in miniature. Its aura of computational precision discourages exactly the questions a citizen most needs to ask: Who built this, on what data, optimizing for which error, accountable to whom? Crawford's larger point is that the scraping and modeling underwriting these systems has become so routine that "few in the AI field even question it" [18] — and if the builders

[18] The Atlas of AI - Power, Politics, and the Planetary Costs

[18] The Atlas of AI - Power, Politics, and the Planetary Costs

have stopped questioning, the burden of scrutiny falls on the public. Transparency is not a courtesy the literate citizen requests; it is leverage she learns to demand.

The cost of not demanding it is no longer hypothetical, and it reaches well past any single domain. When an education chatbot company collapsed, the urgent question — where did the student data go — had no clear answer, because the terms governing that data had never been transparent to the people it described [7]. A whistleblower account of a Los Angeles chatbot deployment described student data misused as the vendor crumbled — a failure of governance that no amount of operational fluency would have caught, but that a culture of transparency-demanding literacy might have forestalled [23]. Investigations into the broader surveillance infrastructure quietly threaded through educational technology reveal how much of this apparatus operates below the threshold of public visibility, harvesting and inferring at a scale users are never invited to evaluate [20]. The literacy that matters here is not the ability to use these systems but the standing to interrogate them — and that standing has to be claimed, because it is never offered.

[7] An Education Chatbot Company Collapsed. Where Did the Student Data - EdSurge

[23] Whistleblower: L.A. Schools' Chatbot Misused Student Data as Tech Co ...

[20] Unmasking EdTech's Surveillance Infrastructure in the Age of AI

Whose Literacy Counts

Every definition of literacy is also a distribution of authority. To say what counts as being literate is to say whose questions are legitimate and whose are noise — and this is where the conceptual muddle becomes a question of power rather than pedagogy.

Notice who currently authors the dominant definitions. The most circulated literacy curricula are produced by the firms that build and sell the tools, and their definition of the literate person is, unsurprisingly, a competent and confident customer [14]. There is a structural conflict of interest in letting the seller define what it means to understand the product. A vendor's ideal of literacy will never include the question "should this system exist at all," because that question is bad for business. The interpretive literacy — the one that preserves the possibility of refusal — has to be authored elsewhere, by publics and the institutions accountable to them rather than to shareholders.

[14] Introduction to AI Literacy - Training | Microsoft Learn

And whose literacy gets *tested* is its own injustice. The false positive, recall, falls disproportionately on non-native English writers, whose prose the detectors misread [2]. The people most likely to be wrongly judged by these systems are precisely the people least likely to have the institutional standing to contest the judgment — a pattern that recurs wherever automated sorting meets unequal power.

[2] AI Cheating in Schools: 2026 Global Trends & Bias Risks

In Göteborg, an automated system for assigning students to schools produced what one careful analysis called a concrete case of algorithmic injustice, the human consequences of a sorting decision the affected families could neither see into nor appeal [19]. When the European Union turned to an American-developed AI to rank its own candidates, the same structure reappeared at the scale of a continent's institutions — the ranked have no window into the ranking [15]. In each case, literacy in the operational sense — knowing how to use the system — belongs to the institution, while literacy in the interpretive sense — the capacity to contest its output — is exactly what the affected population is denied.

This is why the stakes of definition are democratic rather than merely educational. The concern that has surged among parents watching AI enter their children's schooling is, at bottom, a concern about exactly this asymmetry — a sense that decisions are being made by systems they cannot interrogate, on terms they did not set [5]. And it crosses every domain of life, not just schooling. Guidance written for parents navigating AI chatbots and adolescent mental health describes the same predicament in the most intimate register — how to evaluate the trustworthiness of a system speaking to a vulnerable child, when the system's workings are opaque and its incentives unstated [17]. The literate parent, citizen, voter, and patient need the same equipment the wrongly flagged student needs: the probabilistic disposition to distrust the confident number, the source-critical habit to ask how it was made, and the standing to demand an answer.

A 'radical' research effort racing to study artificial intelligence's effect on democracy frames the matter in precisely these civic terms — treating the public's capacity to interpret and contest algorithmic systems not as a consumer skill but as a condition of self-government [1]. That reframing is the whole argument in compressed form. If literacy is fluency, then whoever sells the tools can supply it, and the citizen remains a managed user. If literacy is judgment, then it belongs to the public, and the citizen becomes a person who can say no.

What the False Positive Was Trying to Tell Us

Return, finally, to the student at her kitchen table, falsely accused by a confident machine. We can read that moment two ways, and the reading we choose determines the literacy we build.

The first reading treats the false positive as a bug — a temporary imperfection that a better detector, more training data, a higher

[19] Un cas pratique d'injustice algorithmique : l'attribution automatisée ...

[15] L'UE se tourne vers une IA développée aux États-Unis pour classer ses candidats

[5] AI in the classroom prompts tide of concern from US parents and experts

[17] Parents' Ultimate Guide to AI Chatbots and Mental Health Support

[1] A 'Radical' Lab Races to Study AI's Impact on Democracy

threshold will eventually resolve. This is the reading the operational definition prefers, because it keeps the faith intact: the machine is fundamentally an oracle, currently miscalibrated, soon to be trustworthy. It is also the reading that the empirical record refuses to support, since the sensitivity-specificity tradeoff is not an engineering problem awaiting a fix but a permanent property of classification under uncertainty, and the most rigorous surveys of detector performance keep confirming it [18]. To wait for the infallible detector is to wait for a thing that cannot exist.

The second reading treats the false positive as a feature — not a feature anyone designed, but a revelation the failure provides for free. It tells you that the machine is a probabilistic guesser wearing the costume of a truth-teller. It tells you that every confident output you will ever receive — about a student’s honesty, a video’s authenticity, a candidate’s qualifications, a child’s safety — is an estimate with a hidden margin and an undisclosed cost, and that the cost falls unevenly on the people least able to bear it. Teachers keep using the unreliable tools because uncertainty is uncomfortable [3]; the literate response is to become comfortable with uncertainty, to hold the number at arm’s length, to ask the questions the interface was designed to suppress.

That is the literacy worth wanting, and it is not a skill you can be sold by the people selling the machines. It is the disposition to read a detection score the way you read a poll — with the margin in mind — and to read the institution behind it the way you read any claim of authority: by asking what it is hiding, and demanding that it stop. Janelle Shane’s deflating reminder that everyday AI “is not flawless, not by a long shot” turns out to be the most civically useful sentence in the entire discourse [18]. The false positive was always trying to teach us that. The question is whether we will let ourselves be students of our own machines’ mistakes — or whether we will keep mistaking their confidence for the truth.

References

1. A ‘Radical’ Lab Races to Study AI’s Impact on Democracy
2. AI Cheating in Schools: 2026 Global Trends & Bias Risks

[18] HAI_AI-Index-Report-2024

[3] AI detection tools are unreliable.
Teachers are using them anyway :
NPR

[18] You Look Like a Thing and I
Love You

1. [9] 2. [3] 3. [14] 4. [4] 5. [11] 6. [18] 7. [18] 8. [16] 9. [2] 10. [18]
11. [10] 12. [21] 13. [8] 14. [22] 15. [12] 16. [13] 17. [6] 18. [18] 19. [7]
20. [23] 21. [20] 22. [19] 23. [15] 24. [5] 25. [17] 26. [1]

[9] Colleges pay millions for AI detectors that are flawed - CalMatters
[3] AI detection tools are unreliable.
Teachers are using them anyway :
NPR

[14] Introduction to AI Literacy -
Training | Microsoft Learn

[4] AI for educators - Training |
Microsoft Learn

[11] Frontiers | Prompt engineering as
a new 21st century skill

[18] AI competency framework for
teachers - UNESCO

[18] HAI_AI-Index-Report-2024

[16] More Teachers Are Using AI-
Detection Tools. Here’s Why That
Might Be a ...

[2] AI Cheating in Schools: 2026
Global Trends & Bias Risks

[18] You Look Like a Thing and I
Love You

[10] Deepfakes in the 2024 US Presi-
dential Election

[21] Watch out for false claims of

3. AI detection tools are unreliable. Teachers are using them anyway : NPR
4. AI for educators - Training | Microsoft Learn
5. AI in the classroom prompts tide of concern from US parents and experts
6. AI shows promise in the fight against fake news
7. An Education Chatbot Company Collapsed. Where Did the Student Data - EdSurge
8. Beyond the deepfake hype: AI, democracy, and "the Slovak case"
9. Colleges pay millions for AI detectors that are flawed - CalMatters
10. Deepfakes in the 2024 US Presidential Election
11. Frontiers | Prompt engineering as a new 21st century skill
12. GenAI and misinformation in education: a systematic scoping ... - Springer
13. Inteligencia artificial y desinformación - UNESCO
14. Introduction to AI Literacy - Training | Microsoft Learn
15. L'UE se tourne vers une IA développée aux États-Unis pour classer ses candidats
16. More Teachers Are Using AI-Detection Tools. Here's Why That Might Be a ...
17. Parents' Ultimate Guide to AI Chatbots and Mental Health Support
18. The Atlas of AI - Power, Politics, and the Planetary Costs
19. Un cas pratique d'injustice algorithmique : l'attribution automatisée ...
20. Unmasking EdTech's Surveillance Infrastructure in the Age of AI
21. Watch out for false claims of deepfakes, and actual ... - Brookings
22. What Makes Students (and the Rest of Us) Fall for AI Misinformation?
23. Whistleblower: L.A. Schools' Chatbot Misused Student Data as Tech Co ...