

The Literacy That Runs on a Quota: Why 12 Sessions Won't Save You

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

Somewhere this week a children's guide to artificial intelligence taught its young readers that the way to evaluate a machine's output is to decide whether it is *bueno*—good—or not. That is the whole exercise. You look at what the system produced, you consult your gut, and you render a verdict, the way you might rate a sandwich. It is presented as critical thinking. It is nothing of the kind. It is taste dressed in the vocabulary of judgment, and the distance between the two is the distance this essay is about: the gap between recognizing that something is happening to you and knowing what to do about it.

That gap is now the defining feature of what we call AI literacy. Across the programs, the ministries, the vendor bootcamps and the earnest children's booklets, the field has become expert at teaching people to *notice*—to spot a deepfake, to sense a chatbot's manipulation, to feel that an image is off—while leaving them with no next move. The noticing is real. UNESCO frames the deepfake era as a "crisis of knowing," in which the collapse of our ability to trust what we see corrodes the shared factual ground that democratic life requires [9]. But a crisis of knowing does not get solved by teaching people to feel uncertain more often. It gets solved by giving them power, and power is precisely what the current curriculum withholds.

[9] Deepfakes and the crisis of knowing - UNESCO

This is not an argument that AI literacy is a scam, or that the people building these programs are cynical. Most are not. It is an argument that the concept has been quietly redefined—reduced, actually, in a series of small, plausible-looking steps—until what remains is a literacy that produces better-calibrated consumers and no more capable citizens. The reductions are the story. Let me walk you through them, and then through what we would have to put back.

When Critique Means Saying "Bueno"

Start with that children's guide, because it makes the underlying move visible in miniature. To judge an output as good or bad is to remain entirely inside the frame the system offers you. The machine presents a result; you approve or disapprove; the transaction closes.

What you never do is ask where the result came from, whose labor and whose data it was assembled out of, what it was optimizing for, or who profits when you accept it. Those questions live outside the frame, and the frame is exactly what a taste-based literacy trains you never to leave.

Paulo Freire named this trap half a century ago, and the name still fits with almost insulting precision. In [3], he distinguished the "banking" model of education—in which knowledge is deposited into passive recipients who receive, file, and store it—from a pedagogy of encounter, in which people learn by acting on the world and reflecting on what their action reveals. The banking student is trained to receive verdicts. The encountering student is trained to produce them, and to interrogate the conditions that made the verdict necessary. A curriculum that asks a child whether an AI's answer is *bueno* is banking pedagogy wearing the costume of its opposite. It hands the child a rubber stamp and calls it a scalpel.

[3] *Pedagogia del oprimido*

The deeper problem is that taste, unlike critique, has no leverage on power. Kate Crawford, in [3], warns against what she calls "enchanted determinism"—the habit of treating AI as a marvel to be admired for its innovative method rather than examined for its purpose, a framing that "obscures power and closes off informed public discussion, critical scrutiny, or outright rejection." When you teach people to appraise outputs aesthetically, you are cultivating exactly the enchantment Crawford describes. The system stays a wonder; the child stays a spectator. And a spectator, however discerning, cannot reject anything, because rejection requires standing somewhere other than inside the show.

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

Notice what the taste framing quietly concedes: that the appropriate relationship between a person and an AI system is that of a customer to a product. You sample, you rate, you move on. Crawford's whole argument in the *Atlas* is that this is a category error—that once we "connect AI within these broader structures and social systems, we can escape the notion that artificial intelligence is a purely technical domain," and see it instead as "technical and social practices, institutions and infrastructures, politics and culture" [3]. A literacy worthy of the name would start there, with the institutions and the politics. Instead it starts, and too often ends, with *bueno*.

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

The Quota as Pedagogy

If the children's guide reduces critique to taste, the ministry reduces literacy to attendance. Picture the offer: twelve sessions a year, and—

this is the part that gives the game away—a promise of no pressure. The reassurance is meant to sound humane. It is actually a confession. It tells you that literacy here is conceived as a quantity to be delivered rather than a capacity to be built, a dosage measured in contact hours, with the explicit understanding that nothing much will be asked of you between doses.

An hourly quota is the purest bureaucratic expression of the banking model. Freire's deposits have simply been put on a schedule. The metric of success becomes sessions held, seats filled, completion rates logged—outputs that can be reported upward without anyone ever having to demonstrate that a single participant can now do something they could not do before. This is not a hypothetical failure mode; it is the shape official recommendations already take. The United States' national advisory body, in its own framing document, treats AI literacy largely as a matter of scaling access to instruction across the population [17]—a worthy instinct that nonetheless slides easily into a logistics problem, a question of throughput rather than of what, exactly, is flowing through.

[17] PDF United States of America
RECOMMENDATIONS: Enhancing
AI Literacy for the

Compare the quota to how the field's better instruments describe the thing itself. UNESCO's own literacy materials insist that competencies build in stages, from basic recognition toward the capacity to reason about regulation and social consequence, and that they culminate in judgment about AI's actual uses "for social good" and otherwise across sectors [3]. That is a developmental account: literacy as something that grows, unevenly, through practice and reflection. You cannot deliver it in twelve interchangeable units any more than you can deliver fluency in a language by promising twelve pressure-free evenings a year. The quota model doesn't just under-resource the goal; it misunderstands what kind of goal it is.

[3] UNESCO Think Critically Click
Wisely

And the "no pressure" promise deserves one more look, because it reveals who the program imagines it is serving. Pressure is what you remove when you want a customer to feel comfortable, not when you want a citizen to feel responsible. The framing treats the learner as someone to be kept satisfied rather than someone to be made capable—which is to say it has already decided that the relationship is consumer-service, not civic formation. The MIT Press primer on AI ethics gestures at the alternative: a "more diverse and holistic kind of Bildung," a formation that is radically interdisciplinary in its methods and its topics, the slow cultivation of a whole person's judgment rather than the topping-up of a skills tank [3]. Bildung does not run on a quota. It runs on friction, on being asked to do things you cannot yet do. The ministry has promised to remove precisely the thing that would make its program work.

[3] AI Ethics - The MIT Press Essential
Knowledge series

Recognition Without Response

Here is the pattern that ties the taste guide and the quota program together, and it is the pattern I most want you to carry out of this essay. Nearly every AI literacy effort now teaches recognition and stops there. It teaches you to see the manipulation and abandons you at the moment you might respond to it. This is the field's central, structural failure, and once you notice it you cannot stop seeing it.

Consider the deepfake, the era's signature threat. The instruction is everywhere: look for the uncanny blink, the mismatched shadow, the too-smooth skin. Reporters Without Borders analyzed a hundred deepfakes and warned of a mounting danger specifically to journalists, whose faces and voices are cloned to discredit them [6]. French coverage of the coming municipal elections runs through the same litany—AI, deepfakes, disinformation, democracy under strain [16]. Mexican outlets ask how to halt the spread of AI-generated false news [14]. The recognition curriculum is robust. What follows recognition is a void.

Because suppose you succeed. Suppose you correctly identify a video as synthetic. Then what? You, the individual viewer, have no lever. You cannot compel the platform to remove it, cannot demand a correction reach the people who saw it first, cannot invoke a right that obliges anyone to do anything. The detection tools marketed as a solution are, on inspection, less than they claim—an honest review of the authenticity-verification market finds a landscape of hype where the reliable capabilities are thin and the gap between promise and performance is wide [2]. So the literate viewer is left exactly where the illiterate one is: staring at something they distrust, powerless to act on the distrust.

This is what I mean by recognition without response. The civic remedy exists, but it lives in a domain the literacy curriculum never enters: law, regulation, institutional design, collective action. The person taught to spot a biometric harvesting scheme is almost never taught that a data-protection regime might give them a right to demand deletion, to refuse consent, to file a complaint that carries penalties. The threat is rendered as a matter of personal vigilance; the remedy, which is structural, is left off the syllabus entirely. Carnegie researchers argue that what we actually need is *ongoing, longitudinal monitoring* of how AI shapes political information—an institutional, sustained, collective response, not a moment of individual squinting at a screen [19]. That is the register in which responses live. It is precisely the register the recognition curriculum omits.

[6] Cent deepfakes passés au crible : RSF alerte sur une menace croissante

...

[16] Municipales 2026 : IA, deepfakes et désinformation, la démocratie ...

[14] IA y desinformación: ¿Cómo detener la propagación de ... - MILENIO

[2] AI Content Authenticity Tools: What's Real, What's Hype, and What ...

[19] We Need Ongoing Monitoring of AI and Political Information

The disinformation research bears this out from the supply side. Studies characterizing AI-generated misinformation on social media find it is produced and distributed at a scale and velocity that individual detection cannot match [7]. When experts warn of “bot swarms” engineered to infest social platforms and drown authentic speech, they are describing an industrial phenomenon [12]. You cannot out-notice a swarm. Teaching millions of people to be individually alert to an industrial-scale assault is like handing out umbrellas against a flood and calling it water policy. The Brennan Center’s careful assessment of whether AI fights or fuels election disinformation lands on ambiguity precisely because the outcome depends on institutional choices—on platform rules, on regulation, on enforcement—rather than on how discerning any single voter happens to be [11]. Recognition is the individual’s share. Response is collective. A literacy that delivers only the first has, by design, kept the citizen small.

[7] Characterizing AI-Generated Misinformation on Social Media

[12] Experts warn of threat to democracy from ‘AI bot swarms’ infesting ...

[11] Does AI Fight or Fuel Election Disinformation? - Brennan Center for Justice

The Threats You Can Name But Cannot Answer

The recognition-response gap is easiest to see at the extremes, where the stakes are highest and the missing remedy most glaring. Look at what is happening to teenagers and to the elderly—the two populations onto which AI’s harms are landing hardest, and for whom the literacy on offer is thinnest.

Teenagers have already arrived. A study out of Florida Atlantic University found that sixty percent of American teens have tried AI chatbots and more than one in ten use them daily [1]. This is not a future to be prepared for; it is a present to be survived. And the tools they are using are, by independent assessment, unsafe for the very purpose adolescents most reach for. Common Sense Media found that the major AI chatbots are unsafe for teen mental-health support, offering responses that range from unhelpful to actively dangerous when a young person in distress turns to them [8]. The American Psychological Association has issued a health advisory on AI and adolescent well-being, cataloguing the developmental risks with clinical seriousness [5].

[1] 60% of U.S. Teens Have Tried AI Chatbots, 11.4% Use Them Daily

[8] Common Sense Media Finds Major AI Chatbots Unsafe for Teen Mental ...

[5] Artificial intelligence and adolescent well-being

Now ask what the literacy response to this looks like. Overwhelmingly, it looks like recognition again: teach teens to notice when a chatbot is manipulating them, when a companion app is fostering dependency, when a response feels wrong. But the platforms’ own safety documentation reveals how thin the ground under that recognition is. OpenAI’s help-center page on whether ChatGPT is safe for all ages effectively pushes responsibility back onto the user and the guardian

[20]. The literacy being demanded of the teen is really the abdication of duty by everyone above them. The child is taught to recognize the danger so that the company need not remove it. Recognition, here, is not empowerment; it is liability transfer.

The pattern repeats with grooming and with fraud, where the response gap is not merely an omission but a cruelty. Research on platform-facilitated grooming and AI chatbots documents how conversational systems can be turned into instruments of exploitation, operating in exactly the private, trust-building register that defeats a young person's guard [18]. At the other end of life, guides to AI-powered scams targeting seniors describe voice-cloned grandchildren and synthetic authority figures engineered to bypass a lifetime of accumulated caution [4]. To both populations we say, in effect: be vigilant. But vigilance against a system built specifically to defeat vigilance is not a plan. It is a wish. The response that would matter—platform accountability, enforceable design standards, verification requirements, legal recourse with teeth—belongs to institutions, and it is institutions that the vigilance framing conveniently lets off the hook.

Even the tools sold to parents as a response turn out to reproduce the problem. A review of child-monitoring apps concluded they may need a fundamental rethink, because surveillance of the child substitutes for structural safety and often introduces new harms of its own [15]. This is recognition-without-response escalating into surveillance-as-response: unable to change the system, we are sold ways to watch the victim more closely. The literacy that would actually protect a teenager is not the ability to notice a predatory chatbot. It is membership in a polity that has forbidden the chatbot from being built that way—and knowledge of how to hold that polity to its word.

Literacy as Employability in Someone Else's Cloud

There is a third reduction, and it is the one that most flatters itself as generosity. When a major cloud provider backs a national AI-literacy push, the program is invariably justified by the "talent gap": the economy needs AI-skilled workers, the workers lack the skills, the program closes the gap. Everyone nods. Who could be against closing a gap? But watch the move carefully, because a gap is being closed by a dependency being widened, and the two are not the same thing.

When literacy is defined as employability, its content is set by the employer. You learn the vendor's platform, the vendor's tools, the vendor's certifications—and the vendor, not the public, decides what counts as knowing. This is Freire's banking model again, now

[20] ¿Es ChatGPT seguro para todas las edades? - OpenAI Help Center

[18] Platform-Facilitated Grooming and AI Chatbots

[4] AI Scams Targeting Seniors: 2026 Defense Guide | HCSK

[15] Las apps de monitorización infantil podrían necesitar un ...

with a corporate depositor and a credential at the end. The skills are real and often lucrative; that is not in dispute. What is in dispute is whether acquiring them makes you more powerful or merely more useful, and to whom. A worker fluent in one company's cloud stack is enormously valuable to that company and structurally bound to it. The "gap" that closes is the labor shortage in the vendor's ecosystem. The dependency that widens is the public's, on infrastructure it does not own, cannot audit, and could not replace.

Crawford's *Atlas* is again the sharpest lens here, because it insists on following AI down to its material base—the compute, the data centers, the extracted minerals, the concentrated ownership—and on seeing that this infrastructure is where power actually sits [3]. A literacy that teaches you to operate skillfully atop someone else's infrastructure, while carefully never teaching you who owns it or what that ownership means, has taught you to be a competent tenant and called it citizenship. The enchantment Crawford warns against does its work here too: the training focuses relentlessly on the innovative method, on the marvel of what the tools can do, and never on the prior question of purpose—whose, and why.

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

The civic cost is subtle but real. Employability literacy produces people who can prompt, deploy, and optimize, and who have never once been asked whether the system should exist in the form it does. It answers "how do I use this?" and suppresses "should this be used, by whom, under what rules?" The MIT Press primer's vision of an interdisciplinary *Bildung*—a formation that treats the ethical and political dimensions as internal to competence, so that "to do AI" would simply include the ethics—is the road not taken [3]. The vendor curriculum takes the other fork deliberately, because a workforce that questions the purpose of the infrastructure is worth less to the owner of the infrastructure than one that merely runs it well. That is not a conspiracy. It is just what happens when you let the employer define the literacy: the definition serves the definer.

[3] AI Ethics - The MIT Press Essential Knowledge series

Whose Literacy Counts

Which brings us to the question underneath all three reductions: who gets to define AI literacy, and whose literacy is treated as the real thing? Because the reductions are not random. They all point the same direction—toward a literate consumer and away from a powerful citizen—and that convergence tells you something about the interests doing the defining.

The civic definition, the one being quietly evacuated, would treat

AI literacy as a dimension of democratic participation. The Council of Europe’s work on generative AI and democracy takes exactly this frame, situating the technology’s risks and the public’s competencies within the health of democratic institutions rather than within the labor market or the individual’s media diet [13]. On this account, to be AI-literate is to understand how these systems shape the information environment a self-governing people depends on—and, crucially, to know the institutional levers by which that environment can be governed. It is a definition oriented toward response. It assumes the literate person is a participant in power, not merely a discerning target of it.

[13] IA générative et démocratie

Set that beside the three definitions actually on offer this week. The taste definition makes you a rater. The quota definition makes you an attendee. The employability definition makes you an asset. None of them makes you a participant. And this is where the field’s other silence becomes deafening: the near-total absence of the people least served by every version of it. The literacy discourse is written, overwhelmingly, by and for those with access—the connected, the schooled, the already-online. UNESCO’s own materials note that AI systems still serve only a limited number of spoken languages and continue to encode gendered defaults, which means the “universal” literacy being designed is calibrated to some users and not others before a single lesson is taught [3]. Whose recognition, whose response, whose *bueno*? The question is not rhetorical. A literacy built around the anxieties of the well-resourced will keep producing curricula that protect them and surveil everyone else—the child-monitoring logic scaled to a society.

[3] UNESCO Think Critically Click Wisely

There is also the matter of what the reductions do to trust, which is the substrate all civic literacy runs on. When people are trained to recognize manipulation everywhere and empowered to answer it nowhere, the predictable result is not vigilance but exhaustion—a generalized, corrosive suspicion that everything might be fake and nothing can be done. That is the crisis of knowing UNESCO named, and it is worth being precise about its mechanism: the deepfake’s deepest damage is not that it deceives you into believing a lie but that it teaches you to disbelieve the truth, to shrug at evidence, to retreat from the shared factual world into private cynicism [9]. A recognition-only literacy accelerates exactly this retreat. It hands people the knowledge that they are being manipulated and withholds the means to fight back, which is a recipe for despair, not agency. Even the language guides on countering AI-driven falsehood tend to end where the hard part begins—at the boundary between the individual and the institution [10].

[9] Deepfakes and the crisis of knowing - UNESCO

[10] Deepfakes et désinformation : comment l’IA menace la confiance

The Passage We Keep Refusing to Teach

So what would we have to put back? The answer follows directly from the diagnosis: the entire missing half of the curriculum, the passage from recognition to response.

Freire's word for this passage was *praxis*—reflection and action upon the world in order to transform it, the movement without which, he insisted, education is merely the deposit of inert knowledge [3]. Applied to AI, praxis means that spotting a deepfake is the beginning of a lesson, not its end, and the rest of the lesson is: here is the law that governs synthetic media in your jurisdiction, or here is why no such law exists and here is how one gets made; here is the complaint mechanism and here is what it can and cannot compel; here is the collective body—the regulator, the union, the civic coalition—through which individual recognition becomes structural response. It means teaching the teenager not only that a companion chatbot is engineered to hold their attention but that design standards for such systems are a live policy question they have standing to weigh in on. It means teaching the senior not only that a voice can be cloned but that fraud-liability rules and platform obligations are the actual battlefield, and that they are a constituent in it.

This is harder than the quota. It is harder than the taste rubric and harder than the vendor certification, which is exactly why the field keeps not doing it. Response cannot be delivered in twelve pressure-free sessions, because response is not a quantity you receive; it is a capacity you build by exercising it, ideally alongside other people, against real institutions, with real stakes. It requires the friction the ministry promised to remove and the political content the vendor is structurally incapable of providing. It requires, in short, treating the public as citizens with power to be developed rather than consumers with tastes to be refined or workers with skills to be installed.

The evidence that the current approach is failing is not hidden. It is in the sixty percent of teens already inside systems documented as unsafe for them [1]. It is in the deepfake economy scaling past the reach of any individual's discernment [7]. It is in the calls, from serious institutions, for the sustained collective monitoring that no personal-vigilance curriculum can substitute for [19]. Each of these is a place where recognition has been thoroughly taught and response conspicuously withheld, and the withholding is the wound.

A literacy that runs on a quota was never going to save anyone, because salvation was never the point—throughput was. A literacy that ends at *bueno* was never going to make a citizen, because it was

[3] *Pedagogia del oprimido*

[1] 60% of U.S. Teens Have Tried AI Chatbots, 11.4% Use Them Daily

[7] Characterizing AI-Generated Misinformation on Social Media

[19] We Need Ongoing Monitoring of AI and Political Information

designed to make a rater. And a literacy defined as employability in someone else's cloud will always close the gap it names by widening the dependency it does not. The question worth carrying out of this week's catalog of reductions is the one the reductions are built to keep you from asking: not *can I tell what is fake?*—you increasingly can—but *what, exactly, am I permitted to do about it?* Until the curriculum answers that, it is training spectators for a show in which they were supposed to be the authors. The remedy is not more sessions. It is the one step the whole field keeps refusing to take: from knowing, to acting; from noticing the manipulation, to holding the power that manipulates.

Works cited

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

References

1. 60% of U.S. Teens Have Tried AI Chatbots, 11.4% Use Them Daily
2. AI Content Authenticity Tools: What's Real, What's Hype, and What ...
3. AI Ethics - The MIT Press Essential Knowledge series
4. AI Scams Targeting Seniors: 2026 Defense Guide | HCSK
5. Artificial intelligence and adolescent well-being
6. Cent deepfakes passés au crible : RSF alerte sur une menace croissante ...
7. Characterizing AI-Generated Misinformation on Social Media
8. Common Sense Media Finds Major AI Chatbots Unsafe for Teen Mental ...
9. Deepfakes and the crisis of knowing - UNESCO
10. Deepfakes et désinformation : comment l'IA menace la confiance
11. Does AI Fight or Fuel Election Disinformation? - Brennan Center for Justice

- [9] Deepfakes and the crisis of knowing - UNESCO
- [3] Pedagogia del oprimido
- [3] The Atlas of AI - Power, Politics, and the Planetary Costs
- [17] PDF United States of America RECOMMENDATIONS: Enhancing AI Literacy for the
- [3] UNESCO Think Critically Click Wisely
- [3] AI Ethics - The MIT Press Essential Knowledge series
- [6] Cent deepfakes passés au crible : RSF alerte sur une menace croissante ...
- [16] Municipales 2026 : IA, deepfakes et désinformation, la démocratie ...
- [14] IA y desinformación: ¿Cómo detener la propagación de ... - MILENIO
- [2] AI Content Authenticity Tools: What's Real, What's Hype, and What ...
- [19] We Need Ongoing Monitoring of AI and Political Information
- [7] Characterizing AI-Generated Misinformation on Social Media
- [12] Experts warn of threat to democracy from 'AI bot swarms' infesting ...

12. Experts warn of threat to democracy from 'AI bot swarms' infesting ...
13. IA générative et démocratie
14. IA y desinformación: ¿Cómo detener la propagación de ... - MILENIO
15. Las apps de monitorización infantil podrían necesitar un ...
16. Municipales 2026 : IA, deepfakes et désinformation, la démocratie ...
17. PDF United States of America RECOMMENDATIONS: Enhancing AI Literacy for the
18. Platform-Facilitated Grooming and AI Chatbots
19. We Need Ongoing Monitoring of AI and Political Information
20. ¿Es ChatGPT seguro para todas las edades? - OpenAI Help Center