

# The Fear and the Black Box: How AI Literacy Fails to Interrogate Power

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

Watch what happens when you ask a well-meaning guide to teach you AI literacy in 2026. You will learn to look for the extra finger, the smeared background text, the too-glossy skin, the vowel that lingers a half-beat too long in a synthetic voice. You will be handed a checklist of tells, a forensic vocabulary, a posture of vigilance. What you will almost never learn is who built the system that generated the artifact, what data it was trained on, whose labor was extracted to make it, or under what political and commercial conditions it was released to you in the first place. The literacy on offer is a literacy of surfaces. It teaches you to catch the machine in a lie without ever teaching you to ask who owns the machine.

This is not a small pedagogical quibble. It is the central failure of the way we are currently being taught to live in an AI-saturated world, and it has a shape worth naming precisely. TikTok alone has now labeled more than three billion AI-generated videos, and yet researchers examining those labels find that the labeling scheme systematically misses the content most likely to deceive, because detection is calibrated to obvious synthetic markers rather than to the subtler manipulations that actually move opinion [26]. The label is a comfort object. It signals that the platform is doing something, and it trains the user to believe that the presence or absence of a badge is the meaningful question. It is not.

[26] TikTok Has Labeled 3 Billion AI Videos: Here Is What the Research Says They Miss

The meaningful question is one of power, and the tragedy of contemporary AI literacy is that it has been engineered—sometimes cynically, more often thoughtlessly—to route around that question entirely. What follows is an attempt to map the terrain: to distinguish the competing things people mean when they say “literacy,” to show how the dominant version hollows out the very capacity it claims to build, and to argue that the public is being trained as a class of compliant users when what a democracy needs is a class of citizens capable of oversight.

## *The Two Literacies*

There have always been two things hiding under the word "literacy," and the AI moment has forced them apart. The first is functional: can you operate the tool, recognize its outputs, avoid its most obvious traps? The second is critical: can you interrogate the tool's provenance, its interests, its silent assumptions, and the structures of power it encodes? The distinction is old. UNESCO's own framing of media literacy insists that critical engagement is not an add-on to functional skill but its precondition, and its guidance warns explicitly that AI voice assistants still reproduce gendered stereotypes and serve only a narrow band of the world's languages, meaning that the tool itself carries a politics the user is rarely invited to see [2]. To be literate, in this older and more demanding sense, is to read the medium and not merely the message.

The functional version has won, and it has won because it is legible to institutions, monetizable by vendors, and reassuring to a frightened public. Consider the genre of the deepfake-detection guide, which has metastasized across the web. A representative example walks readers through the visual and acoustic signatures of manipulated media as though disinformation were fundamentally a perception problem to be solved with sharper eyes [8]. The advice is not wrong. It is simply pitched at the wrong altitude. It treats the citizen as an unusually diligent forensic examiner rather than as a political agent with a stake in how synthetic media are produced, funded, and deployed.

The scoping literature on generative AI and misinformation makes the limits of the perception approach unusually clear. A systematic review of the field finds that the relationship between AI-generated content and belief is mediated far less by detectability than by the social and institutional conditions of trust—who is sharing, which networks amplify, what prior commitments the audience brings [12]. You can be perfectly capable of spotting a synthetic artifact and still be perfectly susceptible to the ecosystem that circulates it. The Reuters Institute made the same point about elections with admirable bluntness: the thing to worry about is not the AI system but the humans who wield it, and a literacy trained on artifact-detection leaves the human machinery of manipulation entirely untouched [11].

Here is the conceptual muddle at the heart of the debate: we call both of these things "AI literacy," and by using one word for two incompatible projects we let the weaker one masquerade as the stronger. The functional literacy of tells and badges can be measured, certified, and sold. The critical literacy of power and provenance cannot be reduced to a checklist, which is precisely why it keeps getting quietly

[2] UNESCO Think Critically Click Wisely

[8] Deepfakes et désinformation 2026 : comment les détecter (analyste OSINT ...

[12] Generative AI and misinformation: a scoping review of the role of ...

[11] Generative AI and elections: why you should worry more about humans ...

dropped from the curriculum. When a vendor or a school district announces an AI literacy initiative, it is worth asking which literacy they mean—and the answer, with grim reliability, is the one that produces cooperative users rather than skeptical citizens.

### *Teaching Fear, Not Interrogation*

If the dominant literacy is thin, its affective register is thick with dread. Much of what passes for public education about AI is in fact an education in anxiety, a steady drip of threat scenarios that leaves the learner vigilant, exhausted, and curiously passive. The fraud literature is the purest specimen. Guides to AI voice cloning describe a crisis measured in the billions, walking readers through the mechanics of “vishing” attacks in which a synthetic voice impersonates a family member or an executive to extract money or credentials [7]. Companion pieces interview fraud specialists about the 2026 wave of vocal scams and offer the reader a repertoire of defensive reflexes—call back on a known number, establish a family safe-word, distrust urgency [29].

None of this is bad advice. But notice what it does to the citizen’s relationship to the technology. It positions AI as a natural hazard—a weather system of deception that one must learn to survive—rather than as a set of products built by identifiable companies making identifiable choices about what to release and what to restrain. The safe-word is a coping mechanism, not a form of oversight. It asks nothing of the firms that trained and distributed the voice-cloning models; it asks everything of the potential victim. This is the fear economy of AI literacy: a transfer of responsibility from those who build the systems to those who must endure them.

The workplace variant is more sophisticated but structurally identical. Employer-facing guidance on “information disorder” instructs organizations to inoculate their staff against synthetic content, framing the problem as one of employee credulity to be managed through training and policy [17]. The verification professions feel the same pressure from the other direction. Practitioners of open-source intelligence warn that AI is actively eroding the craft of source verification, flooding the evidentiary environment with plausible fabrications faster than any human process can adjudicate them [16]. A thoughtful assessment from the fact-checking world lays out with precision what AI can and cannot do for verification journalism, and the honest conclusion is that the tools offload some drudgery while introducing new and subtler failure modes that require more judgment, not less [21].

[7] Clonage Vocal IA : La Crise à 1,8 Md\$ qui Brise le KYC

[29] Vishing 2026 : un spécialiste décrypte les arnaques vocales IA

[17] Information disorder in the age of AI: what it means for employers

[16] IA et OSINT : comment l’IA sape la vérification des sources

[21] Lo que la IA puede y no puede hacer por el periodismo de verificación

What unites these documents is an unspoken premise: that the horizon of literacy is defense. You are being taught to protect yourself, your firm, your credulous colleagues. You are not being taught to demand that the systems be built differently, regulated differently, or owned differently. Meredith Broussard’s insistence that we resist the mystification of these systems is instructive here—her argument in *Artificial Unintelligence* is that treating computational systems as inscrutable forces of nature is itself a political act that disarms the public and protects the builders [2]. A fear-based literacy performs exactly this disarmament. It keeps the black box closed and teaches you to flinch at what comes out of it.

[2] Artificial Unintelligence - How Computers Misunderstand

There is a deeper cost. Fear is not a sustainable foundation for civic capacity, because fear seeks relief, and the relief on offer is always more product—better detection tools, premium verification services, the platform’s own labeling scheme. The frightened user does not interrogate the system; the frightened user shops for protection within it. This is how anxiety becomes a business model, and how a literacy that begins in the language of empowerment ends in the reality of dependence.

### *The Technical Frame as Evasion*

Nowhere is the evasion more visible than in the discourse around children, where the stakes are highest and the framing most revealingly narrow. When a new report found that Google’s AI search features posed what its authors called an unacceptable risk to children—surfacing dangerous content, fabricating authoritative-sounding falsehoods, collapsing the distinction between a vetted source and a statistical guess—the coverage registered the alarm but rendered the problem as a matter of technical safety to be patched [1]. The implicit remedy is better filters. The unasked question is why a single corporation was permitted to insert an unreliable generative layer between hundreds of millions of children and the world’s information in the first place, and what recourse the public has when it does.

[1] ‘It’s deeply disturbing.’ What a new report says about risks Google’s AI search features pose to kids

The pattern repeats across the child-safety literature with almost mechanical consistency. The Grok sexual deepfake scandal, in which a major AI system was used to generate non-consensual intimate imagery, became a story about content moderation and platform response rather than about the corporate decision to ship an image model with those capabilities largely unconstrained [13]. Reporting on the deepfake crisis in schools documented a phenomenon far worse than administrators had assumed—students weaponizing synthetic

[13] Grok sexual deepfake scandal - Wikipedia

imagery against classmates—and framed it, again, as a safety and disciplinary challenge rather than as the predictable downstream consequence of releasing these tools into the world without accountability [20]. In each case the technical frame does a specific kind of work: it converts a question of power into a question of engineering.

Consider the surveillance response, which is where the technical frame turns actively pernicious. When schools deploy AI monitoring systems like Gaggle to watch students, the tools carry documented risks of false alarms, privacy violation, and the chilling of expression—and yet they are marketed and adopted as child-protection technology, a safety measure that happens to require total surveillance [22]. The mental-health domain shows the same move. France’s data protection authority surveyed young people’s use of conversational AI for emotional support and found a landscape of dependency and risk that regulators are only beginning to grasp [15]. A widely circulated argument holds that teenagers need guardrails rather than bans on mental-health chatbots—a reasonable-sounding position that nonetheless takes the existence and desirability of the chatbots as given, and quietly relocates the entire debate onto the terrain of the vendors’ preferred solution [23].

The consumer-guide genre completes the picture. Guides promising “safe AI for teens” proliferate, offering parents settings to toggle and conversations to have, all of which presuppose that safety is a configuration problem soluble at the level of the individual household [14]. This is the technical frame in its domestic form: a literacy that hands you a dashboard and calls it agency. Kate Crawford’s argument in *The Atlas of AI* cuts against exactly this reduction—her insistence is that artificial intelligence is never merely a technical artifact but “a manifestation of highly organized capital backed by vast systems of extraction and logistics,” a claim that the child-safety discourse works overtime to keep out of view [2]. To frame the harm to children as technical is to promise that it can be fixed without touching the organization of capital that produced it. That promise is false, and a literacy built on it is worse than useless, because it teaches the public to accept the terms of a bargain it was never allowed to negotiate.

### *Models as National Assets*

There is a dimension of AI literacy that the skills-based curriculum cannot even see, because it lies entirely outside the frame of individual competence: the fact that the most capable models are not civic resources but strategic assets, subject to geopolitical control, national

[20] La crisis de deepfakes en las escuelas es mucho peor de lo que ...

[22] Programas de IA para monitorear a estudiantes tienen riesgos de ...

[15] IA conversationnelle et santé mentale des jeunes - CNIL

[23] Teens need guardrails, not bans, for mental health chatbots | STAT

[14] Guía completa: IA segura para adolescentes (2026) | HolaNolis

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

security logic, and corporate discretion. You can be maximally literate in the functional sense—fluent with every prompt, alert to every tell—and still discover one morning that the tool you built your practice around is no longer available to you, because a government or a company decided so.

The security discourse has already arrived at this framing, though it arrives brandishing threat rather than analysis. When a member of Congress warned that China had deployed “digital twins” of every American lawmaker—AI models trained to simulate and predict the behavior of political figures—the story was pitched as a national-security alarm [5]. Strip away the alarm and what remains is a structural fact: advanced AI is being treated by states as an instrument of statecraft, a capability to be hoarded, weaponized, and denied to adversaries. The citizen’s relationship to these systems is therefore mediated not only by markets but by borders and by the shifting definitions of who counts as a threat.

This matters for literacy in a way that the detection-and-defense curriculum cannot accommodate. If model availability is conditional and political, then a literacy that prepares you only to use the tools has prepared you for a world that may not persist. The more fundamental competence is understanding the dependency itself—recognizing that the generative layer now woven through search, work, and public life is controlled by a handful of firms operating under national jurisdictions, and that this control is exercised in ways the public neither authorizes nor sees.

The political speech question sharpens the point. When Meta’s Oversight Board assessed whether large language models suppress political expression, it found genuine cause for concern: that the models mediating an increasing share of public discourse make consequential and largely invisible decisions about what speech to permit, dampen, or refuse [4]. These are not neutral tools awaiting your skilled use. They are governance systems, enforcing content policies that no electorate ratified, in domains—political speech, civic information—that a democracy is supposed to keep under public control. A literacy adequate to this reality would have to teach the citizen to see the model as a site of governance and to ask by what authority it governs.

The electoral application makes the danger concrete. A serious policy analysis of chatbot voting advice concludes that these systems could both help and harm voters and therefore demand regulation calibrated to their influence—an admission that we have already begun outsourcing civic judgment to opaque corporate models whose training, incentives, and errors the voter cannot inspect [6]. When a citizen

[5] Cammack says China has deployed ‘digital twins’ of every lawmaker

[4] Are LLMs Stifling Political Speech? An Assessment of How ...

[6] Chatbot Voting Advice Could Help and Harm Voters. Let’s Regulate Accordingly.

asks a chatbot how to vote, the black box is no longer a curiosity; it is a participant in the democratic process, one whose interests are undisclosed. Noam Chomsky’s account of how manufactured consent operates through control of the informational commons is uncomfortably apposite here—the mechanism has simply migrated from editorial gatekeeping to algorithmic mediation, and the public has even less visibility into the new gatekeepers than it had into the old. A literacy that cannot name this migration cannot defend against it.

### *The Deskilling of Judgment*

There is a further erosion, quieter than fraud and less dramatic than geopolitics, but corrosive precisely because it feels like progress. As professionals across fields fold generative tools into their daily judgment, they risk atrophying the very capacities that made their judgment worth having. The Spanish-language business press has given this its sharpest name—the “deskilling” challenge, in which AI threatens to atrophy the skills of an organization’s best employees as they outsource the difficult cognitive work that once kept those skills sharp [10]. The phenomenon is not new—every tool reshapes the skills of its user—but the scale and the domain are. What is being outsourced now is judgment itself.

The mechanism deserves precise description, because it connects the deskilling problem back to the power problem that runs through this whole terrain. Research on what has been called “the generative illusion” documents how ChatGPT-like systems produce outputs so fluent, so confident, so formally competent that users systematically overestimate their reliability and underinvest in verifying them [24]. The fluency is the trap. A system that produced obvious garbage would keep the user’s critical faculties engaged; a system that produces plausible prose lulls them to sleep. And the very design choices that make the output fluent—the ones that maximize engagement and perceived competence—are made by the firms that build the models, in service of adoption, not in service of the user’s continued capacity to think.

This is why the deskilling story is a literacy story rather than merely a labor story. When judgment migrates into the model, the ability to interrogate the model becomes both more necessary and less available, because the skills that would enable the interrogation are the ones being eroded. NewsGuard’s early demonstration that ChatGPT could function as what it called the next great misinformation superspreader—generating fluent, authoritative falsehoods

[10] El desafío del “deskilling”: cuando la inteligencia artificial amenaza con atrofiar habilidades de los mejores empleados

[24] The generative illusion: how ChatGPT-like AI tools could ... - Frontiers

on demand—illustrates the endgame: a public that has outsourced its judgment to systems designed to sound trustworthy regardless of whether they are [25]. The deskilled user cannot tell the difference, because telling the difference is exactly the skill that has withered.

The institutional response has largely made things worse by mistaking the problem. A widely shared argument holds that AI is not killing education but merely exposing what was already broken about it—a bracing thesis that nonetheless treats the tools as a diagnostic revelation rather than as an active force reshaping cognition [3]. Meanwhile the vendors have moved to occupy the pedagogical space directly. Anthropic’s launch of a dedicated teaching product is framed as support for learning, but it also represents something more consequential: the entity that builds the black box positioning itself as the authority on how the public should be taught to relate to it [18]. When the maker of the model also authors the literacy curriculum, the curriculum’s silence on the model’s power is not an oversight. It is the product working as designed.

The French cultural-policy discourse offers a rare counterweight, insisting that generative AI’s arrival in education and culture demands a defense of critical spirit—an *esprit critique* developed collectively rather than delivered by the tools themselves [19]. But notice how marginal such arguments remain against the momentum of adoption. The critical spirit is expensive, slow, and unmonetizable; the fluent tool is cheap, fast, and everywhere. Left to the market, the deskilling will proceed not because anyone chose it but because no one with power had an interest in preventing it.

### *Whose Literacy Counts*

Every literacy encodes a theory of who the learner is and what they are being made fit for. The functional literacy that dominates our moment makes the learner into a user—someone to be equipped for safe and productive consumption of systems designed elsewhere by others. A critical literacy adequate to democratic life would make the learner into something else: a citizen with a legitimate claim to know how the systems that govern their information, their work, and their political life actually operate. The gap between these two figures is the gap between being managed and being empowered, and it is worth asking who benefits from keeping the public on the wrong side of it.

The transparency debate is where this question becomes concrete, and where the limits of the current settlement are most exposed. Legal scholars examining algorithmic transparency have begun to map the

[25] The Next Great Misinformation Superspreader: How ChatGPT ... - NewsGuard

[3] AI Isn’t Killing Education. It’s Exposing What Was Already Broken.

[18] Introducing Claude for Teachers

[19] L’IA générative, l’éducation et la culture

boundaries of the "right to know"—the uncomfortable finding that even where such rights exist, they are hedged by trade-secret protections, technical opacity, and the sheer complexity of the systems, such that the right to an explanation often delivers an explanation no one can use [27]. A public that cannot understand the explanation it is legally owed is not, in any meaningful sense, being given oversight. It is being given the appearance of oversight, which is worse, because the appearance forecloses the demand.

Where genuine civic infrastructure has been built, it has come from public pressure rather than corporate initiative, and it remains fragile. The Spanish investigative organization Civio's campaign for an algorithm registry—a public record of the automated systems governments use to make decisions about citizens—represents a concrete vision of what democratic AI literacy might actually require: not skilled users but institutional transparency, enforced by law, legible to the public [28]. Notice that this is a literacy located not in the individual's head but in the structure of the state—a recognition that no amount of personal skill can substitute for a right to know exercised collectively.

This reframing matters because it redistributes responsibility. The dominant literacy places the entire burden of coping on the individual: learn the tells, set the parental controls, establish the safe-word, verify the output. A democratic literacy asks instead what the individual is owed—transparency, accountability, recourse—and treats the cultivation of that demand as the core civic skill. Even the genuine goods of these systems have to be understood in this register. AI can be a real instrument of accessibility for people whose brains work differently, opening capacities that were previously foreclosed [30]. But even the emancipatory case is written by the vendor, and a public equipped only to receive the benefit—rather than to ask on what terms, at what cost to privacy, under whose control the benefit is delivered—has been handed a gift it cannot govern.

The proliferation of AI-generated disinformation across the media environment, now a recurrent rather than exceptional problem, closes the loop [9]. The labeling schemes proposed as remedies return us to where we began: to the badge, the tell, the surface. A public trained to look for labels has been trained to accept a world in which the fundamental questions—who produces this flood, who profits from it, who could stop it—are settled without them. The UNESCO framework for teachers is right that AI-generated content is inherently "less trustworthy" because of its stochastic nature, but trustworthiness is not finally a technical property to be assessed [2]. It is a relationship of accountability between the public and the powerful, and it is that relationship, not the artifact, that a real literacy would teach the public to

[27] Transparencia algorítmica: límites del derecho a saber

[28] Un registro de algoritmos, el primer paso para una IA transparente en ...

[30] When your brain works differently, AI isn't a luxury—it's accessibility | Artificial Intelligence

[9] Des contenus générés par IA sources de plus en plus récurrentes de ...

[2] AI competency framework for teachers - UNESCO

examine.

### *The Compliant User and the Absent Citizen*

The through-line of this terrain is a substitution so smooth that most people never notice it happening: the citizen has been replaced by the user, and the replacement has been sold as empowerment. The user is taught to detect, to configure, to defend, to adapt. Every one of these verbs accepts the system as given and asks only how the individual might survive within it. The citizen would be taught different verbs—to interrogate, to demand, to regulate, to refuse—and every one of those accepts nothing as given, treating the system instead as a human construction answerable to public authority. We are producing users at industrial scale and citizens almost not at all.

This is not because the citizen’s literacy is impossible to teach. It is because the parties with the resources to teach at scale—the platforms, the vendors, the security industry that profits from fear—have no interest in a public that asks who owns the machine. The literacy they fund and distribute is precisely calibrated to leave the black box closed. When Anthropic writes the teaching materials, when TikTok designs the labels, when the voice-cloning guide sells you a safe-word, the curriculum’s silence about power is not a gap to be filled. It is the point. Broussard’s warning against the mystification of computational systems names the mechanism exactly: opacity is not an accident of complexity but a distribution of advantage, and it advantages the builders every time [2].

What would it mean to teach the other literacy? It would mean beginning not with the artifact but with the arrangement—treating every encounter with an AI system as an occasion to ask who built it, who profits, who is accountable, and who was never consulted. It would mean, as Civio’s algorithm registry insists, building the public’s capacity to know into institutions rather than leaving it to the diligence of frightened individuals [28]. It would mean taking seriously the Oversight Board’s finding that these models now govern speech, and insisting that governance requires consent [4]. It would mean refusing the fear economy and its endless upsell of protection, and replacing the question “how do I stay safe?” with the harder and more democratic question “by what right does this system operate on me?”

The fear and the black box are made for each other. The fear keeps the public looking at the surface—the extra finger, the suspicious badge, the urgent voice on the phone—while the box stays sealed and the power inside it goes uninterrogated. A public that can spot

[2] Artificial Unintelligence - How Computers Misunderstand

[28] Un registro de algoritmos, el primer paso para una IA transparente en ...

[4] Are LLMs Stifling Political Speech? An Assessment of How ...

a deepfake but cannot name the interests it serves has been given a flashlight and told it is a key. The most capable AI systems in the world are being treated, openly now, as national and corporate assets whose availability is conditional and whose logic is opaque, and the literacy on offer prepares the public for none of it. Until we insist on the other literacy—the one that opens the box rather than teaching us to fear its contents—we will keep producing exactly the citizenry that concentrated power prefers: skilled, vigilant, anxious, and profoundly unable to ask the only question that matters.

---

1. [26] 2. [2] 3. [8] 4. [12] 5. [11] 6. [7] 7. [29] 8. [17] 9. [16] 10. [21]  
 11. [2] 12. [1] 13. [13] 14. [20] 15. [22] 16. [15] 17. [23] 18. [14] 19. [2]  
 20. [5] 21. [4] 22. [6] 23. [10] 24. [24] 25. [25] 26. [3] 27. [18] 28. [19]  
 29. [27] 30. [28] 31. [30] 32. [9] 33. [2]

## References

1. 'It's deeply disturbing.' What a new report says about risks Google's AI search features pose to kids
2. AI competency framework for teachers - UNESCO
3. AI Isn't Killing Education. It's Exposing What Was Already Broken.
4. Are LLMs Stifling Political Speech? An Assessment of How ...
5. Cammack says China has deployed 'digital twins' of every lawmaker
6. Chatbot Voting Advice Could Help and Harm Voters. Let's Regulate Accordingly.
7. Clonage Vocal IA : La Crise à 1,8 Md\$ qui Brise le KYC
8. Deepfakes et désinformation 2026 : comment les détecter (analyste OSINT ...
9. Des contenus générés par IA sources de plus en plus récurrentes de ...
10. El desafío del "deskilling": cuando la inteligencia artificial amenaza con atrofiar habilidades de los mejores empleados
11. Generative AI and elections: why you should worry more about humans ...

[26] TikTok Has Labeled 3 Billion AI Videos: Here Is What the Research Says They Miss

[2] UNESCO Think Critically Click Wisely

[8] Deepfakes et désinformation 2026 : comment les détecter (analyste OSINT ...

[12] Generative AI and misinformation: a scoping review of the role of ...

[11] Generative AI and elections: why you should worry more about humans ...

[7] Clonage Vocal IA : La Crise à 1,8 Md\$ qui Brise le KYC

[29] Vishing 2026 : un spécialiste décrypte les arnaques vocales IA

[17] Information disorder in the age of AI: what it means for employers

[16] IA et OSINT : comment l'IA sape la vérification des sources

[21] Lo que la IA puede y no puede hacer por el periodismo de verificación

[2] Artificial Unintelligence - How Computers Misunderstand

[1] 'It's deeply disturbing.' What a new report says about risks Google's AI search features pose to kids

[13] Grok sexual deepfake scandal - Wikipedia

12. Generative AI and misinformation: a scoping review of the role of ...
13. Grok sexual deepfake scandal - Wikipedia
14. Guía completa: IA segura para adolescentes (2026) | HolaNolis
15. IA conversationnelle et santé mentale des jeunes - CNIL
16. IA et OSINT : comment l'IA sape la vérification des sources
17. Information disorder in the age of AI: what it means for employers
18. Introducing Claude for Teachers
19. L'IA générative, l'éducation et la culture
20. La crisis de deepfakes en las escuelas es mucho peor de lo que ...
21. Lo que la IA puede y no puede hacer por el periodismo de verificación
22. Programas de IA para monitorear a estudiantes tienen riesgos de ...
23. Teens need guardrails, not bans, for mental health chatbots | STAT
24. The generative illusion: how ChatGPT-like AI tools could ... - Frontiers
25. The Next Great Misinformation Superspreader: How ChatGPT ... - NewsGuard
26. TikTok Has Labeled 3 Billion AI Videos: Here Is What the Research Says They Miss
27. Transparencia algorítmica: límites del derecho a saber
28. Un registro de algoritmos, el primer paso para una IA transparente en ...
29. Vishing 2026 : un spécialiste décrypte les arnaques vocales IA
30. When your brain works differently, AI isn't a luxury—it's accessibility | Artificial Intelligence