

The Proctored Mind: How Universities Chose Surveillance Over Teaching

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

Somewhere in the middle of this week’s higher-education reporting sits a number that ought to have stopped an administrator cold: a reported rise of sixty-seven percent in “suspected AI fraud” cases, brandished in press releases and dashboards as if it were a public-health statistic, a measure of an epidemic sweeping through the residence halls. What is remarkable is not the number itself but the silence around it — the absence of the one question a curious institution might have asked before it reached for the alarm. Nobody, in the telling, wondered whether the assessments producing those flagged cases were themselves the problem. The number was treated as evidence about students. It is, in fact, evidence about the university.

That reflex — to read a crisis of trust as a crime wave rather than a design failure — is the through-line of everything worth noticing in the discourse right now. When the *Los Angeles Times* describes trust between colleges and their students as visibly “decaying” under the pressure of AI cheating accusations, it is describing a mutual condition, but the institutional response has been asymmetrical [20]. The university has decided that the erosion is happening *to* it, perpetrated *by* the people it enrolls, and the remedy it has purchased — detection software, remote proctoring, keystroke analytics, webcam gaze-tracking — is a remedy addressed entirely to the suspected party. This is the quiet transformation I want to trace here: the slow conversion of the trust architecture that once made teaching possible into a surveillance architecture, undertaken not through some grand ideological decision but through a series of small procurements made whenever a crisis appeared and the pedagogical work seemed too hard.

The largest empirical study of undergraduate AI use to date, out of Berkeley, is worth holding in mind as a corrective before we go any further, because it complicates the panic in exactly the way the panic refuses to complicate itself [18]. What the researchers found was not a homogeneous population of cheaters but a stratified one — disparities in who has access to the tools, disparities in who uses them for what, and a picture of student behavior far more legible as a response to conditions than as a moral collapse. The institution that reads that study and hears only “cheating is up” has already decided what kind

[20] Trust in colleges decaying over AI cheating - PressReader

[18] The largest study of AI use by undergrads is in, revealing disparities in access and in cheating

of problem it wants to have.

The Cheating Panic Is a Design Confession

Consider the take-home exam, that quiet fixture of the modern syllabus. For decades it functioned on an implicit bargain: we cannot watch you, so we trust you, and in exchange we ask questions whose answers you could look up but whose reasoning you cannot fake. Generative AI did not break that bargain so much as expose how thin it had always been. When a *New York Times* investigation concludes that student cheating is becoming "impossible to detect" in the age of AI, the honest reading is not that students have become more devious but that a whole genre of assessment was quietly outsourcing its integrity to the friction of the library and the limits of the copy-paste [11]. Remove the friction and the assessment reveals that it was never really testing what it claimed.

The take-home exam and the integrity-violation spike are not two stories; they are one story told from two ends. An assessment that can be completed by a language model without the student learning anything was, before the language model arrived, an assessment that could be completed by a diligent copyist, a well-organized study group, or a student with an older sibling's paper in a drawer. The model simply industrialized what was always possible and made it cheap. The institutional confession embedded in the spike is that a great deal of what universities have been grading was never a measure of thinking in the first place — it was a measure of compliance with a format. This is why the most useful writing this week comes from those willing to say the quiet part: that the era of generative AI forces us, as one French account of the debate puts it bluntly, to reconsider whether the exam itself needs to be rethought rather than merely defended [22].

The most vivid case of the year should have settled the argument. When the biggest AI cheating scandal in the Ivy League erupted at Brown University — a mass-scale integrity breach that, in a detail almost too grim to credit, unfolded in the aftermath of a campus shooting — the framing offered by the coverage was overwhelmingly about student malfeasance and institutional betrayal [17]. But an event of that scale is not a story about individual bad actors. A cheating incident that can involve a substantial fraction of a cohort is a story about an assessment regime that made cheating both easy and, in the students' moral calculus, defensible. When misconduct becomes normative, the norm being violated has already lost its authority. The

[11] Las trampas de los estudiantes se están volviendo imposibles de detectar en la era de la IA

[22] Utiliser ChatGPT, est-ce tricher : à l'ère des IA génératives, faut-il repenser les examens

[17] The biggest Ivy league AI cheating ever happened after a mass shooting

scandal is the syllabus.

None of this is an argument that students never cheat, or that the pressures on them excuse dishonesty. It is an argument about where the institution chooses to point its attention and its budget. Even the vendors' own admission-of-scale numbers betray the logic. When Latin America's largest university, the UNAM, was forced to investigate mass irregularities in its admission examination, the institutional instinct was again to treat the breach as a security failure to be surveilled away rather than a design problem to be engineered out [15]. The pattern is remarkably stable across systems and languages: a crisis appears, and the reflex is to buy suspicion rather than to build better questions.

[15] Por primera vez, la universidad más grande de América Latina hizo público un escándalo en su examen de admisión

When a Library Guide Becomes the Whole Strategy

If the assessment crisis reveals what universities do under pressure, the AI-literacy vacuum reveals what they do when nobody is watching. Here the confession is one of omission. Across a striking number of institutions, the entirety of the official response to generative AI — the whole considered pedagogical posture of a research university — has been delegated to a library research guide: a webpage, maintained by a subject librarian of goodwill and no institutional power, listing a few tools, a few cautions, a link to the honor code. This is not a criticism of librarians, who have done more thoughtful work on this question than most faculty senates. It is a criticism of the decision to let their guide *be* the strategy.

Housing AI literacy in a library guide rather than in the curriculum is a policy decision, and it is a decision with visible consequences for who learns what. A curriculum is mandatory, sequenced, assessed, and resourced; a guide is optional, static, and encountered only by the student who already knows to look. When the weekly higher-education research briefings catalogue institutional responses, the gap between the volume of activity and the depth of it is conspicuous — a great deal of guidance, a great deal of policy, and very little of it embedded where learning actually happens [14]. The literacy that students most need — the capacity to judge when a model's fluent output is wrong, to understand what it cannot do, to use it as an instrument rather than an oracle — is precisely the literacy that cannot be acquired from a bulleted list of "acceptable uses."

[14] PDF Weekly AI in Higher Education Report

The equity consequences of this omission are not abstract. When AI literacy is left to be picked up informally, it reproduces the disparities the Berkeley study documented: the students who arrive already

fluent pull further ahead, and those who arrive uncertain are left to improvise, often in the direction of the very misuse the institution then punishes. The point is sharpened by the research on assessment and language equity, which shows how the line between legitimate support and illegitimate substitution falls differently on students writing in a second language, who use these tools to do things — smoothing grammar, checking idiom — that a monolingual peer never needed help with [9]. An institution without a taught literacy has no coherent way to draw that line, so it draws it with a detector, which draws it wrongly.

[9] GenAI Assessment and Language Equity: Drawing the Line between Support and Substitution

The absence is most damaging for the students the institution claims to serve most carefully. The policy work on students with disabilities makes the case that AI tools are, for many of them, not a shortcut but an accommodation — a text-to-speech scaffold, a drafting aid, a way around a barrier that has nothing to do with the competence being assessed [13]. A university that has not thought its way to a curriculum-level position on AI literacy cannot distinguish accommodation from cheating, and its detection regime, applied uniformly, will systematically misread the former as the latter. The library guide, well-meaning and toothless, cannot carry that weight. It was never meant to.

[13] PDF Prioritizing Students With Disabilities in AI Policy

What makes this an *unrecognized surrender* — the phrase is deliberate — is that the vacuum does not stay empty. Where the institution declines to build a curriculum, the vendors arrive with a product, and the curriculum quietly becomes whatever the tool permits. The business schools have at least been candid enough to war-game the outcome. An unusually honest piece from the AACSB asks the question most strategy documents avoid — what happens if the AI strategy actually *succeeds*, if the efficiencies are real and the adoption complete — and finds the answer unsettling, because a fully successful automation strategy hollows out the very apprenticeship that professional education exists to provide [23]. The uncomfortable implication is that the institutions with no strategy and the institutions with a triumphant one may arrive at the same destination: a curriculum defined by what the software does, chosen by nobody in particular.

[23] What If Our AI Strategy Succeeds?

The Carceral Premise Nobody Will Name

Nowhere is the surrender to vendors clearer than in the proctoring market, and nowhere is the animating premise more nakedly displayed. The single most clarifying document of the week is *WIRED*'s account of a student's campaign against the exam-surveillance software Proc-

torio — a piece that, almost incidentally, exposes the assumption on which the entire industry is built [19]. The software watches the test-taker through a webcam, tracks eye movement, flags “abnormal” behavior, and demands the student prove continuous innocence for the duration of the exam. Its defenders treat its failures — the facial-recognition systems that cannot see dark-skinned faces, the algorithms that flag a student’s disability as suspicious fidgeting — as bugs to be patched. The student’s insight, and *WIRED*’s, is that these are not bugs. They are the honest expression of the design’s foundational assumption: that the student is a suspect, and the exam is an interrogation.

This is the carceral premise, and it deserves to be named plainly because the discourse works so hard to keep it unnamed. A systematic review of online proctoring systems, published through the ERIC database, catalogues the documented harms with academic restraint — the false positives, the privacy intrusions, the anxiety, the demonstrable bias — and yet even it frames these as problems of implementation to be optimized [12]. The vendors go further, insisting the whole apparatus is a matter of fairness misunderstood: the trade publications of the proctoring industry argue that what is needed is simply *clearer rules*, as though the problem with a surveillance regime were insufficient documentation of its terms [1]. The premise itself — that continuous suspicion is an acceptable condition of learning — is never on the table.

The harms are not hypothetical, and they are not evenly distributed. The phenomenon that the mental-health literature has taken to calling “flagxiety” — the corrosive dread of being falsely flagged by a detector one cannot see and cannot appeal — is now documented as a genuine campus mental-health problem, a tax levied on every honest student for the sake of catching a few dishonest ones [6]. The detectors themselves do not work well enough to justify that tax: the technical reviews are increasingly blunt that AI-detection tools in the academy produce error rates incompatible with the disciplinary weight placed on them [7]. And the errors fall hardest on the students already most vulnerable — second-language writers, whose more formulaic prose reads to a detector as machine-generated, a pattern now generating an actual wave of litigation [5].

Then there is the failure mode that no amount of clearer rules can patch: the catastrophic one. When an AI-supervised remote examination went so badly wrong that fifty-eight thousand students had to retake it, the event demonstrated something the vendors’ incrementalism cannot absorb [3]. An institution that has outsourced its assessment to a surveillance vendor has also outsourced the single point of failure.

[19] This Student Is Taking On ‘Biased’ Exam Software | WIRED

[12] PDF A Systematic-Narrative Review of Online Proctoring Systems and a Case Study

[1] Academic Integrity in the Age of AI: Why Clear Proctoring Rules Matter

[6] Detección de IA “Flagxiety”: la crisis de salud mental del campus

[7] Detectores de IA en la academia: ¿funcionan?

[5] Demandas por detección de IA 2026: guía para escritores ESL

[3] An AI-supervised remote exam went so badly that 58,000 students must retake it

The fragility is not a temporary condition of an immature market; it is what dependence looks like. This is the dynamic Shoshana Zuboff describes when she warns that learning in society has been “hijacked by surveillance capitalism,” and that in the absence of institutions strong enough to tether it, we are thrown back on the market form of the surveillance companies “in this most decisive of contests” over the division of learning [16]. The university that buys a proctoring platform is not buying a tool; it is ceding a jurisdiction.

[16] The Age of Surveillance Capitalism

Kate Crawford makes the lineage explicit. The methods now filtering into classrooms — the behavioral tracking, the anomaly detection, the automated accusation — are the same methods once “reserved for extralegal espionage,” and she notes with precision how welfare-fraud systems learned to “track anomalous data patterns in order to cut people off” and accuse them, before those tools filtered down “to classrooms, police stations, workplaces, and unemployment offices” [16]. The proctoring webcam is not a neutral descendant of the exam invigilator. It is a descendant of the fraud-detection apparatus, and it carries that apparatus’s founding assumption — that the population under observation is a population of probable cheats — directly into the seminar room. The Pulitzer Center’s reporting on campus AI surveillance frames the question with appropriate bluntness: is this security, or is it a privacy invasion dressed as security? [21]. The institutions buying the systems have mostly declined to answer.

[16] The Atlas of AI

[21] Using AI on Campuses: Security Surveillance or Privacy Invasion?

The Socratic Bargain Requires Risk

There is a deeper pedagogical loss underneath the procurement decisions, and it is easy to miss because it is a loss of something the surveillance apparatus was never designed to touch. The oldest technique in the teacher’s repertoire — the Socratic question, the deliberate exposure of the student to the productive discomfort of not-yet-knowing — depends on precisely the thing the current institutional posture is engineered to eliminate: risk. The Socratic method works because the student is permitted to be wrong in public, to venture an answer that might collapse under the next question, to sit in uncertainty long enough for understanding to form. Remove the risk of being wrong and you remove the mechanism. What is left is theater.

The trouble is that the entire architecture of AI-mediated learning, as currently procured, is a machinery for removing risk. Consider the strongest evidence in AI’s favor. A rigorous randomized controlled trial published in *Nature* found that AI tutoring can outperform in-class active learning on measured outcomes — a genuinely

striking result, and one the skeptics should not wave away [2]. But the Brookings synthesis of the tutoring research adds the necessary complication: the gains are real and also narrow, concentrated in well-defined domains with clear right answers, and they say little about the messier, riskier work of learning to think in an ill-structured field [24]. An AI tutor is superb at removing the friction of not-knowing quickly. That is exactly what makes it a poor instrument for the kind of learning that requires the friction to remain.

The psychological literature is beginning to name the cost. The American Psychological Association's reporting on how AI is reshaping human skills and thinking gathers the emerging evidence that habitual cognitive offloading — letting the system carry the effortful part — erodes the very capacities the effort was building [10]. A student who never sits in the discomfort of a hard problem, because a fluent answer is always one prompt away, does not develop the tolerance for difficulty that is, arguably, the whole point of an education. This is the Socratic bargain inverted: the technique is hollowed out precisely when the risk is removed, and yet risk-removal is what the institution, in its procurement decisions, keeps buying — smoother tutoring, faster answers, tighter surveillance, less friction, less exposure, less of the productive discomfort in which learning lives.

There is a version of AI-supported pedagogy that runs the other way, that uses the tool to *increase* the stakes rather than launder them — the experiential-learning research points toward it, showing gains when AI is used to put students into consequential simulations rather than to shortcut them [8]. But that version requires a curriculum designed by educators who have thought hard about where risk belongs. It cannot be bought from a vendor, and it cannot be delegated to a library guide, which returns us to the surrender: the institution that will not do the pedagogical work will default to the tools that promise to make the work unnecessary, and those tools are, almost without exception, instruments for the elimination of exactly the risk that teaching requires.

The Word the Discourse Cannot Say

Read across the week's reporting and one absence becomes deafening. In all the talk of detection and integrity and policy and surveillance, there is almost no talk of *partnership* — no framing in which the student is a collaborator in the work of learning rather than a suspect to be monitored or a liability to be managed. The institution has a vocabulary rich in the language of enforcement and impoverished in

[2] AI tutoring outperforms in-class active learning: an RCT

[24] What the research shows about generative AI in tutoring

[10] How AI is reshaping human skills and thinking

[8] Effects of AI guided experiential learning on market ...

the language of trust, and the vocabulary is not accidental. It reflects a settled choice about what kind of relationship the university believes it has with the people it educates.

The irony is that the collaborative framing is available, and it comes from an unexpected quarter. The most constructive institutional advice on what to do when a student submits AI-generated work comes not from a university but from OpenAI itself, whose guidance for educators counsels conversation, clarity of expectation, and treating the incident as a teaching moment rather than a tribunal [4]. One should be appropriately wary of a vendor playing the reasonable adult — the company selling the shovels has every reason to counsel calm about the gold rush. But the fact that the most pedagogically humane advice in circulation is coming from the AI company rather than from the academy is itself the indictment. The institution has abandoned a posture, and someone with a product to sell has been happy to occupy it.

Meredith Broussard, whose skepticism about AI is bracing precisely because it is not reflexive, insists that the way forward requires being “more honest and realistic about what goes into technology and what it takes to keep technology working,” and holds out the possibility of “a path forward that uses technology to support both democracy and human dignity” [16]. Human dignity is the missing term in the proctoring debate. A surveillance regime is, by construction, a regime that treats the observed party as an object of suspicion rather than a subject worthy of trust, and no amount of clearer rules or better facial recognition can convert an architecture of suspicion into an architecture of dignity. Those are different buildings. You cannot renovate one into the other.

What a partnership framing would actually require is not softness but rigor of a different kind — the harder rigor of assessment design that makes cheating pointless because the task demands the thinking, of literacy taught in the curriculum rather than posted in a guide, of a relationship in which the student’s use of a powerful tool is a subject of open instruction rather than covert detection. The Iberoamerican and Québécois policy frameworks now circulating gesture at this more considered posture, treating AI’s arrival as an occasion for deliberate institutional design rather than defensive reaction — but they remain, for now, documents rather than practices, aspiration rather than architecture [14]. The gap between the thoughtful framework and the deployed detector is the gap between what universities say about themselves and what their procurement reveals.

[4] Comment les éducateurs peuvent-ils réagir lorsque des élèves présentent comme leur un contenu généré par l’IA

[16] Artificial Unintelligence

[14] PDF Weekly AI in Higher Education Report

What the Architecture Remembers

The buildings we choose to construct outlast the crises that prompted them. A surveillance apparatus, once installed, does not dismantle itself when the panic subsides; it acquires a budget line, a vendor relationship, an administrative constituency, and a logic that seeks new applications. The proctoring platform bought to catch pandemic-era remote cheats becomes the analytics platform monitoring engagement, becomes the risk-scoring system flagging students for intervention, becomes — following exactly the lineage Crawford traces — a permanent fixture whose founding assumption of suspicion quietly recolors the whole institution’s understanding of the people inside it [16]. This is why the choice being made this week matters beyond this week. The institution is not merely responding to a cheating crisis; it is deciding what kind of place it will be once the crisis is forgotten.

[16] The Atlas of AI

The evidence assembled here does not describe a technology imposing itself on unwilling institutions. It describes institutions making choices — to read a design failure as a crime wave, to leave the curriculum empty and let vendors fill it, to treat the carceral premise of proctoring as a fixable flaw rather than a disqualifying one, to buy the removal of the very risk that learning requires. Each choice was locally reasonable, made under pressure, by people with too much to do. Together they amount to the surrender the stakes named at the outset: the conversion of education’s trust architecture into a surveillance architecture, undertaken without anyone quite deciding to undertake it.

The most damning fact remains the simplest one. When trust between colleges and their students decays, as the reporting plainly shows it has, an institution genuinely committed to teaching would ask what it might do to rebuild the trust [20]. The institution we actually have has mostly asked how it might better monitor the people it no longer trusts. Zuboff’s warning is that in the contest over who controls the division of learning, surveillance capital has “gathered the power and asserted the authority to supply all the answers” [16]. The university’s task was to be the institution that supplied better questions — questions worth answering, asked in a relationship worth having. It is not too late to choose the questions. But the proctored mind is a hard thing to un-proctor, and the architecture, once built, remembers what it was built to assume.

[20] Trust in colleges decaying over AI cheating - PressReader

[16] The Age of Surveillance Capitalism

References

1. Academic Integrity in the Age of AI: Why Clear Proctoring Rules Matter
2. AI tutoring outperforms in-class active learning: an RCT
3. An AI-supervised remote exam went so badly that 58,000 students must retake it
4. Comment les éducateurs peuvent-ils réagir lorsque des élèves présentent comme leur un contenu généré par l'IA
5. Demandas por detección de IA 2026: guía para escritores ESL
6. Detección de IA "Flagxiety": la crisis de salud mental del campus
7. Detectores de IA en la academia: ¿funcionan?
8. Effects of AI guided experiential learning on market ...
9. GenAI Assessment and Language Equity: Drawing the Line between Support and Substitution
10. How AI is reshaping human skills and thinking
11. Las trampas de los estudiantes se están volviendo imposibles de detectar en la era de la IA
12. PDF A Systematic-Narrative Review of Online Proctoring Systems and a Case Study
13. PDF Prioritizing Students With Disabilities in AI Policy
14. PDF Weekly AI in Higher Education Report
15. Por primera vez, la universidad más grande de América Latina hizo público un escándalo en su examen de admisión
16. The Age of Surveillance Capitalism
17. The biggest Ivy league AI cheating ever happened after a mass shooting
18. The largest study of AI use by undergrads is in, revealing disparities in access and in cheating
19. This Student Is Taking On 'Biased' Exam Software | WIRED
20. Trust in colleges decaying over AI cheating - PressReader
21. Using AI on Campuses: Security Surveillance or Privacy Invasion?

22. Utiliser ChatGPT, est-ce tricher : à l'ère des IA génératives, faut-il repenser les examens
23. What If Our AI Strategy Succeeds?
24. What the research shows about generative AI in tutoring