

The Algorithm That Grades You Before You Apply: AI Sorting in Credit, Work, and Justice

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

There is a peculiar violence in being judged before you have said a word. A loan officer who once had to meet your eyes, a hiring manager who at least had to read your name, a parole board that had to sit in a room with you — all of these have been quietly replaced by a score that arrives before you do. The application is no longer the beginning of the process; it is the ratification of a verdict already reached. By the time you click “submit,” the algorithm has already decided what kind of person you are, and it has done so using a portrait of you that you never sat for.

This is the machinery worth watching this week, because the reporting across credit, employment, and criminal justice converges on a single, under-examined fact: algorithmic sorting is not malfunctioning. It is functioning exactly as designed. When Amsterdam spent years and considerable public money trying to build a demonstrably *fair* welfare-fraud algorithm and abandoned the effort after discovering that fairness in one dimension simply migrated the harm into another, the lesson was not that the engineers were careless [17]. The lesson was that the sorting itself — the ranking of citizens by predicted risk — was the harm, and no amount of technical adjustment could subtract it. The tool worked. That was the problem.

[17] Inside Amsterdam’s high-stakes experiment to create fair welfare AI

What I want to trace here is a question of power: who gets to build these grading systems, who gets graded, who narrates the harm when it surfaces, and — most importantly — who is conspicuously absent from the conversation about whether the grading should happen at all. Because the striking feature of this week’s discourse is not that AI harms are hidden. They are extravagantly documented. The striking feature is that the documentation flows in one direction, from the graded to the sympathetic observer, while the power to redesign flows in another, and the two currents almost never meet.

The Grade That Precedes You

Start with credit, because credit is where the logic is oldest and therefore least questioned. The premise of a credit score is that your financial past predicts your financial future, but the contemporary version has quietly dissolved the boundary of what counts as your "financial past." Scoring systems now ingest digital footprints — browsing behavior, social-network structure, the phone you use, the hour you fill out a form — and convert the residue of ordinary life into a risk signal. Shoshana Zuboff named the underlying business model in [24]: the unilateral claim over private human experience as free raw material, rendered into behavioral data and traded for predictions. Credit scoring is that logic made actuarial. Your consent is neither sought nor meaningful, because the surveillance happens upstream of the moment you would think to withhold it.

[24] The Age of Surveillance Capitalism

Kate Crawford's argument in [24] sharpens the point: classification is not a neutral description of the world but an act of power, and once embedded in working infrastructure it becomes "relatively invisible without losing any of its power." A credit algorithm does not merely observe that some borrowers are riskier than others; it *constitutes* the categories of risk, and it does so in ways that reliably track existing hierarchies. Ruha Benjamin, in [24], describes how systems that "seem to objectively rank people on the basis of merit" invoke efficiency and progress as the lingua franca of innovation — a rhetoric that launders sorting into something that sounds like fairness. The Chinese social-credit apparatus she cites, which promises to "forge a public opinion environment where keeping trust is glorious," is not an aberration from the Western model but its honest face.

[24] The Atlas of AI

[24] Race After Technology

The bias in these systems is not evenly distributed, and the reporting confirms whose life gets read as risk. Analysis of AI systems across Latin America documents how models trained on data from wealthier, whiter, more male populations systematically misread everyone else — reproducing gender, racial, and xenophobic bias not as an occasional error but as a structural default [10]. A broader survey of algorithmic bias reaches the same conclusion from a different angle: these tools reflect and amplify the ideas already dominant in their training data, which is to say the ideas of those with the power to generate data in the first place [18]. The person who has been surveilled the most — the frequent applicant, the poor, the migrant — is the person the system claims to know best, and knows worst.

[10] Género, racismo y xenofobia: así son los sesgos de la Inteligencia

[18] Inteligencia artificial y sesgos: cómo una IA puede reflejar ideas

Bias as Bug, Harm as Feature

The employment domain is where the sleight of hand becomes easiest to catch, because here the industry has settled on a comforting story: bias is a bug, and bugs can be audited away. The story has the great advantage of being partly true and entirely convenient. It is true that hiring algorithms encode discrimination. It is convenient because it locates the problem in the code rather than in the decision to automate labor screening at all.

The canonical illustration predates this week but haunts all of it. When Amazon built a résumé-screening tool, engineers eventually discovered that among the features most strongly correlated with a candidate’s predicted success were being named Jared and having played high-school lacrosse — a finding recounted in [24]. The model had faithfully learned what Amazon’s past hiring looked like and proposed to reproduce it forever. What deserves attention is the response: the engineers tried to *delete* the offending variables, as though the bias were a stain to be scrubbed rather than the signal the system was built to detect. This is the audit paradigm in miniature. It treats the discriminatory output as a defect in an otherwise sound project, when the project itself — predicting worth from historical pattern — guarantees that yesterday’s exclusions become tomorrow’s criteria.

This week the audit paradigm collided with the courts. The litigation against Workday, whose applicant-screening software an aggrieved plaintiff alleges rejected him from hundreds of jobs on the basis of protected characteristics, has been allowed to proceed as a collective action potentially reaching millions of applicants [1]. The case matters less for its likely damages than for what it exposes: candidates rejected by an algorithm had no way to know they had been screened by one, no way to see the criteria, and no way to contest the outcome — a case study in the machinery of un-appealable rejection [12]. The Spanish-language coverage frames the stakes bluntly: a court has effectively told millions of applicants that a piece of software may have discriminated against them, years after the fact and with no remedy that reaches most of them [6].

Meanwhile the same logic is arriving at the other end of the employment relationship. Meta faces what is described as the first lawsuit alleging that AI-driven layoff selection discriminated against workers, with a claims deadline pressing employees to act before they even understand how the system ranked them [21]. The workplace is being wrapped in a second layer of scoring — continuous profiling and surveillance of employees, whose legal and ethical hazards observers are only beginning to map [25]. The score no longer merely decides

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

[1] AI Hiring Discrimination: How Algorithms Reject Millions of Qualified

[12] How AI Bias Locked Out Millions of Job Seekers (A Case Study on Mobley

[6] Discriminación de la IA en la contratación: la demanda contra Workday y

[21] Meta Faces First AI Layoff Discrimination Suit as July 22 Deadline Looms

[25] The Legal and Ethical Minefield of A.I.-Driven Employee Surveillance

whether you get in; it decides whether you stay, and on what terms.

Regulators have noticed, and their response is instructive precisely because it accepts the industry’s frame. California has moved to tighten scrutiny of AI hiring tools amid documented racial bias, treating the problem as one of insufficient testing and disclosure [3]. The advice offered to employers across Latin America follows the same script: audit your algorithms, remove the bias, hire more fairly [14]. Note what this frame forecloses. It asks whether the screening machine is biased; it never asks whether the screening machine should hold the power it holds. There is a further, quieter danger the reporting names — the *monoculture* of the algorithm, in which every employer licenses the same handful of scoring engines, so that a candidate rejected by one is effectively rejected by all, with no diversity of judgment left in the labor market to appeal to [22]. A single biased human is a misfortune. A single biased algorithm deployed across an entire economy is a structure.

[3] California Tightens Scrutiny of AI Hiring Tools Amid Reports of Racial Bias

[14] IA en reclutamiento: cómo auditar los sesgos algorítmicos y contratar

[22] Monocultivo algorítmico en contratación y sesgo sistémico

The Thirty-Nine Percent That Speaks Loudest

Step back from the individual cases and look at the shape of the discourse itself, because the shape is telling. Across the week’s reporting on AI and society, the dominant register is the documentation of ethical failure — bias exposed, surveillance revealed, discrimination litigated. Nearly two in five of these stories are, at their core, failure narratives. This is usually read as a sign of health: the watchdogs are watching, the harms are surfacing, accountability is coming. I want to suggest a less flattering reading.

The over-representation of ethical-failure stories reveals a discourse that has become fluent in diagnosis and mute on redesign. We have an enormous and growing archive of *what went wrong* — the Amazon tool, the Workday screen, the Amsterdam welfare model, the Latin American bias audits — and a conspicuously thin archive of anyone with power proposing that the sorting be dismantled rather than corrected. The authors of [24] concede a structural reason this asymmetry persists: purely legal and regulatory oversight “doesn’t scale,” because any system relying on human attention “cannot possibly keep up with the volume and velocity of algorithmic decision-making.” The harms are produced at machine speed; the accountability is produced at human speed, one lawsuit and one exposé at a time. The gap between them is not closing. It is the permanent operating condition.

[24] The Ethical Algorithm

This asymmetry has an author, and the author is power. Consider who narrates each harm. The bias in hiring algorithms is documented

by journalists, plaintiffs’ attorneys, and academic auditors — people with standing to describe the wound but no authority to close the factory that makes it. The people with that authority — the vendors licensing the scoring engines, the enterprises deploying them, the agencies procuring them — appear in the discourse mainly as defendants issuing statements or as sponsors of the very fairness audits that certify their systems as improved. The result is a conversation in which the harmed speak volubly about their injuries and the powerful speak sparingly about their intentions, and the structural question — should this scoring exist? — is asked by neither, because the harmed lack the platform and the powerful lack the incentive.

The consequence is a discourse optimized to feel critical while remaining safe. Every documented failure implicitly promises that the *next* version will be fairer, which is to say that the failure narrative, far from threatening the expansion of algorithmic sorting, underwrites it. Each exposé becomes a specification document for the next generation of tools. This is the trap Benjamin identifies in [24]: the reform that fixes the surface bias while deepening the underlying capture. When we celebrate that an algorithm has been “de-biased,” we have accepted that the algorithm will continue to sort — we have merely negotiated the terms of our own ranking.

[24] Race After Technology

Actuarial Justice and Its Missing Question

Nowhere is the refusal to ask the structural question more consequential than in criminal justice, where the score does not decide a loan or a job but liberty itself. The field of algorithmic risk assessment has, for over a decade, promised to reduce human bias in bail, sentencing, and parole decisions by replacing the prejudiced judge with the neutral calculation. The empirical record is, at best, ambiguous. A careful review of the evidence finds that risk-assessment algorithms have not delivered the bias reduction they promised and may entrench the disparities they were meant to cure, because the data they learn from — arrests, convictions, prior contact with police — is itself the sediment of discriminatory policing [2].

What deserves scrutiny is how the scholarship responds to this record. The dominant move is to stay inside the paradigm — to refine the risk instrument, recalibrate its thresholds, adjust its fairness metrics — rather than to ask whether predicting the dangerousness of human beings from statistical aggregates is a legitimate function of a justice system at all. This is what critics call actuarial justice: the treatment of defendants not as individuals who did or did not do a

[2] Algorithms Were Supposed to Reduce Bias in Criminal Justice—Do They?

specific act but as members of a risk class, priced like insurance. The research literature debates *which* fairness definition to optimize, tacitly conceding that the tools will be used. The prior question — whether the state should grade its citizens’ future crimes before they are committed — is left to philosophers and abolitionists, voices with little purchase on procurement decisions.

And procurement continues apace, increasingly through the front door of surveillance rather than the back door of sentencing. Facial recognition has become the most visible edge of the carceral algorithm, and the pattern of its deployment maps the power asymmetry precisely. When public outcry forced the Milwaukee police to pause facial-recognition adoption, the pause was provisional — “for now,” as the reporting carefully noted, a concession negotiated under pressure rather than a principled limit [23]. Where cities have imposed outright bans, police have simply found other routes to the same capability, accessing the banned technology through partnerships, third parties, and jurisdictional workarounds [8]. The technology adapts faster than the prohibition. A new generation of AI now lets police identify people by gait, body shape, and other biometric signatures precisely to *skirt* facial-recognition bans — the regulation aimed at faces, the surveillance simply moved to bodies [11].

The stakes are clearest where the graded population has the least recourse. ICE agents now wield an expanded toolkit of facial recognition and phone-tracking spyware to identify and pursue immigrants, a population defined by its exclusion from the political community that might otherwise constrain the tools [15]. Governments across the world have turned the same recognition systems on protesters, converting the act of public dissent into a biometric record [13]. Crawford’s warning in [24] — that welfare and enforcement systems now deploy “capture that was once reserved for extralegal espionage” against ordinary people, tracking anomalous patterns to cut them off and accuse them of fraud — is no longer a warning. It is a description of the operating environment.

The global legal patchwork underscores that this is a question of power, not of technical readiness. Some jurisdictions ban facial recognition; others build national systems around it; the difference correlates less with the maturity of the technology than with the distribution of political voice among those it targets [9]. Where the graded can vote and litigate, the tool meets friction. Where they cannot — the migrant, the protester, the poor — it advances unimpeded. The score is not applied evenly because power is not distributed evenly, and the algorithm is, in the end, a faithful instrument of that distribution.

[23] Public outcry over facial recognition technology leads Milwaukee police

[8] Estas ciudades prohíben la tecnología de reconocimiento facial. La

[11] How a new type of AI is helping police skirt facial recognition bans

[15] ICE agents have new tools to track and ID people : NPR

[13] How governments use facial recognition for protest surveillance - Rest

[24] The Atlas of AI

[9] Facial Recognition Laws Worldwide: Who Bans It, Who Builds It

The Ghost Workers Behind the Grade

Every grading algorithm rests on a foundation of human labor that the discourse of AI almost never grades in return. The models that sort credit applicants, screen résumés, and flag suspects are trained on data that must be cleaned, labeled, and moderated by people — and those people constitute the largest structural silence in the entire conversation. They are the workers who make the machine appear autonomous by disappearing themselves.

The reporting on this labor is harrowing and, revealingly, comes from the margins of the coverage rather than its center. In India, women are paid to watch hours of abusive and violent content so that AI systems can learn to recognize and filter it, absorbing the psychological damage as an occupational externality — “in the end,” as one described it, “you feel blank” [16]. Across the poorer economies of the world, hundreds of thousands of workers perform the invisible data labeling that makes the whole apparatus function, for wages that would be illegal in the countries where the resulting tools are sold [20]. And there is a bleak recursion in the arrangement that one investigation names precisely: this is the labor that helps AI eliminate labor, the human effort spent training the systems that will render human effort redundant [7].

This is where the power analysis turns from metaphor to geography. The extraction that Crawford anatomizes in [24] — the planetary supply chain of minerals, energy, and labor beneath every “intelligent” system — has a linguistic and cultural dimension as well. Critical work on data colonialism documents how indigenous languages are scraped, modeled, and commodified by AI systems built elsewhere, without the consent or benefit of the communities that speak them, reproducing the colonial pattern of extracting raw material from the periphery to enrich the center [4]. The people whose lives and languages are the substrate of these systems are precisely the people with no seat at the table where the systems are designed. The continent-scale version of this argument is being made about Africa, where analysts warn that AI, far from enabling industrial development, may lock the region into dependency — a consumer of imported models and a supplier of raw data, never an author of the tools that grade it [27].

The silence extends to who counts as a legible subject at all. GLAAD’s analysis of AI’s impact on LGBTQ people documents how systems trained on normative data misread, erase, or actively endanger those who fall outside the categories the training data anticipated — a reminder that the grade is calibrated to a default person, and everyone who deviates from that default pays a tax in misrecognition [26]. Even

[16] In the end, you feel blank’: India’s female workers watching hours of abusive content to train AI

[20] Los cientos de miles de trabajadores en países pobres que hacen - BBC

[7] El trabajo que ayuda a la IA a acabar con el trabajo

[24] The Atlas of AI

[4] Data colonialism and indigenous languages in AI: a critical review of

[27] Why AI can hold back Africa’s industrialisation and what to

[26] Understanding LGBTQ Impacts Across AI – 2026 AI Report

the algorithm’s tone carries the freight: research shows conversational AI reproduces gender stereotypes in the very manner of its speech, reinforcing hierarchies so subtly that the companies deploying it do not perceive the bias they are transmitting [19]. The people best positioned to notice these harms — the misgendered, the mistranslated, the ghost-worked — are the people least represented among those who build and govern the tools. That is not a coincidence to be corrected. It is the arrangement the tools exist to preserve.

[19] La IA también te habla con sesgo: cuando el tono algorítmico refuerza estereotipos de género sin que la empresa lo perciba

Governance That Tinkers at the Edges

If the harms are so well documented, why does the sorting keep expanding? The answer lies in the character of the governance response, which has settled, almost everywhere, on tinkering rather than questioning. The dominant regulatory instinct is to make the algorithm fairer, more transparent, more accountable — to improve the grading rather than to interrogate whether grading on this scale should be permitted. This is not a failure of governance. It is governance performing its actual function, which is to legitimize the expansion by attaching to it the vocabulary of oversight.

Watch the vocabulary do its work. The language of “fairness,” “bias mitigation,” and “responsible AI” is deployed most enthusiastically by the institutions with the most to gain from the systems’ continued deployment. Benjamin’s account in [24] of how ranking systems invoke “efficiency and progress as the lingua franca of innovation” applies exactly to the governance layer: fairness has become the lingua franca by which extractive scoring is made palatable. When Amsterdam demonstrated, through an unusually honest and well-resourced attempt, that a fair welfare algorithm could not in fact be built — that improving one fairness metric degraded another, with no configuration that harmed no one — the finding should have been treated as a boundary condition on the entire enterprise [17]. Instead it is filed as a hard problem awaiting a better solution.

[24] Race After Technology

[17] Inside Amsterdam’s high-stakes experiment to create fair welfare AI

The Danish case shows what the tinkering conceals. Amnesty International’s investigation found that Denmark’s AI-powered welfare system functions as an engine of mass surveillance, flagging marginalized groups — the disabled, migrants, the poor — for fraud investigation on the basis of proxies that amount to discrimination by another name [5]. This is a wealthy, rights-conscious democracy with robust data-protection law. The governance did not prevent the harm; it accompanied it. The system is legal, audited, and discriminatory all at once, which is precisely the outcome a fairness-focused governance regime is

[5] Dinamarca: El sistema de bienestar social basado en la inteligencia

designed to produce: not the absence of harm but its regularization.

At the geopolitical scale, the governance conversation reveals its true stakes as a contest over who controls the sorting infrastructure rather than whether sorting should occur. When Xi Jinping calls for greater global efforts to regulate AI while criticizing the United States for throttling the technology's spread, the framing is control, not restraint — a struggle over which power sets the standards that everyone else will be graded by [28]. No major power in this contest proposes that the citizens of the world be exempted from algorithmic scoring. They propose only that their own scoring standards prevail. The authors of [24] are right that regulation "doesn't scale" against machine-speed decision-making — but the deeper problem is that regulation, as currently practiced, does not aim to scale against the sorting. It aims to manage the sorting's optics.

[28] Xi pide más esfuerzos globales para regular la IA y critica a EEUU por frenar la tecnología

[24] The Ethical Algorithm

What the Score Cannot See

Return, finally, to the person judged before they apply. The credit applicant whose social graph priced her risk before she asked for a loan. The job seeker rejected by a screening engine he never knew was running, with no diversity of judgment left in a monocultured market to appeal to [22]. The defendant priced as a risk class. The immigrant identified by his gait after the ban on his face. The Indian content moderator who went blank so that a filter could learn. What unites them is not that they were harmed by a broken system but that they were correctly processed by a working one — sorted, scored, and slotted into hierarchies the algorithm inherited and now enforces at scale.

[22] Monocultivo algorítmico en contratación y sesgo sistémico

The discourse this week is loud with their injuries and quiet about their design. It documents the bias exhaustively and questions the sorting almost never, because the power to document belongs to observers and the power to redesign belongs to the sorted's sorters, and those two constituencies do not overlap. Crawford's insight in [24] — that classification embedded in infrastructure becomes invisible without losing its power — is the whole story compressed: the more normal the grading becomes, the less it is noticed, and the less it is noticed, the more freely it expands. The fairness audit, the transparency report, the provisional ban "for now" — each is a way of noticing the machine while leaving it running.

[24] The Atlas of AI

A genuinely reader-serving conclusion cannot end on the reassurance that better tools are coming, because the better tool is the trap. The question worth carrying out of this week is not whether the al-

gorithm that grades you can be made fairer. It is whether a society should grant any actor — corporate or state — the standing to grade you before you speak, to price your future from your past, and to call the result progress. The people best placed to answer that question are the ones the systems were built to sort, and they are the ones the conversation has not yet learned to hear.

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

References

1. AI Hiring Discrimination: How Algorithms Reject Millions of Qualified
2. Algorithms Were Supposed to Reduce Bias in Criminal Justice—Do They?
3. California Tightens Scrutiny of AI Hiring Tools Amid Reports of Racial Bias
4. Data colonialism and indigenous languages in AI: a critical review of
5. Dinamarca: El sistema de bienestar social basado en la inteligencia
6. Discriminación de la IA en la contratación: la demanda contra Workday y
7. El trabajo que ayuda a la IA a acabar con el trabajo
8. Estas ciudades prohíben la tecnología de reconocimiento facial. La
9. Facial Recognition Laws Worldwide: Who Bans It, Who Builds It
10. Género, racismo y xenofobia: así son los sesgos de la Inteligencia
11. How a new type of AI is helping police skirt facial recognition bans
12. How AI Bias Locked Out Millions of Job Seekers (A Case Study on Mobley
13. How governments use facial recognition for protest surveillance - Rest

- [17] Inside Amsterdam's high-stakes experiment to create fair welfare AI
- [24] The Age of Surveillance Capitalism
- [24] The Atlas of AI
- [24] Race After Technology
- [10] Género, racismo y xenofobia: así son los sesgos de la Inteligencia
- [18] Inteligencia artificial y sesgos: cómo una IA puede reflejar ideas
- [24] You Look Like a Thing and I Love You
- [1] AI Hiring Discrimination: How Algorithms Reject Millions of Qualified
- [12] How AI Bias Locked Out Millions of Job Seekers (A Case Study on Mobley
- [6] Discriminación de la IA en la contratación: la demanda contra Workday y
- [21] Meta Faces First AI Layoff Discrimination Suit as July 22 Deadline Looms
- [25] The Legal and Ethical Minefield of A.I.-Driven Employee Surveillance
- [3] California Tightens Scrutiny of AI Hiring Tools Amid Reports of Racial Bias

14. IA en reclutamiento: cómo auditar los sesgos algorítmicos y contratar
15. ICE agents have new tools to track and ID people : NPR
16. In the end, you feel blank': India's female workers watching hours of abusive content to train AI
17. Inside Amsterdam's high-stakes experiment to create fair welfare AI
18. Inteligencia artificial y sesgos: cómo una IA puede reflejar ideas
19. La IA también te habla con sesgo: cuando el tono algorítmico refuerza estereotipos de género sin que la empresa lo perciba
20. Los cientos de miles de trabajadores en países pobres que hacen - BBC
21. Meta Faces First AI Layoff Discrimination Suit as July 22 Deadline Looms
22. Monocultivo algorítmico en contratación y sesgo sistémico
23. Public outcry over facial recognition technology leads Milwaukee police
24. The Ethical Algorithm
25. The Legal and Ethical Minefield of A.I.-Driven Employee Surveillance
26. Understanding LGBTQ Impacts Across AI – 2026 AI Report
27. Why AI can hold back Africa's industrialisation and what to
28. Xi pide más esfuerzos globales para regular la IA y critica a EEUU por frenar la tecnología