

GUBMENT SPECIAL EDITION WHITEPAPER NO. SE1 · SERIES GBMT-SE1 · FILED
2026-08-04

AI: we're regulating a quantity nobody measures

A feasibility assessment of the leading "govern AI" theories — built around a measurement gap that makes both sides of the harm debate unfalsifiable, a safety regime no one can audit, and a power bill in the tens of billions that nobody files under AI policy.

READ

90 SECONDS

10 MINUTES

FULL RECORD

CONTENTS / RECORD

+

ABSTRACT

This inquiry set out to name the binding constraints on governing artificial intelligence and score the leading architectures against them. It fired its own kill condition on the first pass. No federal instrument measures how often AI actually decides anything consequential — not in hiring, not in lending, not in housing; in healthcare, the one federal rule that touches it shows which algorithms exist to identified staff inside the deploying organisation — not to the public, and not how often they decide, and a pending rule would strip even that. Every claim that AI is causing widespread discrimination, and every claim that it isn't, is therefore unfalsifiable at population scale. Three findings follow from that gap rather than around it. First, the harms that *have* reached a court or a regulator are the unglamorous ones — discrimination settlements, \$893M in losses on fraud complaints that reference AI, 400,000-plus AI-linked child-exploitation reports — while the catastrophic-risk scenarios dominating the discourse have produced essentially zero adjudicated cases. Second, the safety commitments meant to govern those scenarios cannot be verified by anyone *in the United States*: the federal evaluation body was renamed, narrowed, and left with voluntary participation, and no US instrument — federal or state, enacted or pending — grants an outside party standing access to a lab's weights, training data, or compute logs. That sentence originally ran without the words "in the United States," and the unqualified version is **withdrawn**: the EU AI Act gives the Commission a compulsory, fine-backed power to obtain a general-purpose model — source code included — and to appoint independent experts to evaluate it, in force since two days before this filing published. Third, the largest measurable public cost of the AI build-out is a utility bill: Georgia alone certified \$50–60 billion in ratepayer exposure, and the mid-Atlantic grid's own market monitor attributes \$29.4 billion in capacity charges to data centers — sums that rival total federal AI spending and are almost never discussed as AI policy. A nine-architecture scorecard, red-teamed before publication (a rankings

table that didn't match its own inputs, caught and rebuilt; two cells corrected), lands — after an independent blind re-score corrected eleven cells — on state-law primacy and disclosure mandates co-leading, frontier-lab licensure sharing the floor with the voluntary status quo, the comprehensive federal statute mid-pack as an all-neutral row, and one newly separated result: voluntary safety commitments score worse than doing nothing at all, because published assurance that independent assessment grades unreliable is documented theater, while absence at least makes no false claim. **A Phase 2 steelman has since shown that the first three of those results each turn on one cell** — and that the reason given here for the third, that the implemented EU evidence supports nothing above neutral, is wrong on the statute. Only the fourth held under every variant tested.

- 1 · The quantity nobody measures
- 2 · The harms that actually reach a court
- 3 · Nobody can check the safety homework
- 4 · Your power bill is AI policy
- 5 · The jobs numbers have two parents
- 6 · Washington blocks; it doesn't build
- 7 · The coalition math
- 8 · What the old regimes teach
- 9 · The scorecard, corrected
- × · The honesty box

PART 1

The quantity nobody measures



0

FEDERAL INSTRUMENTS MEASURING AI IN HIRING DECISIONS

21.5%

OF US BUSINESSES USING AI AT ALL (CENSUS, JUL 2026)

0

IN LENDING

0

IN HOUSING

This filing named a test in advance: if no national instrument measures AI deployment in consequential decisions, that finding becomes the headline. It does not exist. The EEOC withdrew its AI hiring guidance in January 2025 and collects no AI-use data; the CFPB withdrew sixty-seven guidance documents including its AI circulars in May 2025, and the national mortgage dataset has no AI-use field. New York City's hiring-AI audit law is the closest working instrument in the country — one city, and a December 2025 state comptroller audit found its enforcement ineffective, with three-quarters of complaint calls misrouted. Healthcare is the partial exception that proves the rule: federal certification requires health-IT vendors to make source-attribute descriptions of their predictive algorithms available to a limited set of identified users — not to the public — at the 96% of hospitals reporting a certified EHR, but not how often those algorithms drive a decision, and a pending rule proposes removing the requirement entirely.

phase0-findings §1 · EEOC/CFPB withdrawal record (90 FR 20084) · 45 CFR 170.315(b)(11), ONC HTI-1 (89 FR 1192) and proposed HTI-5 (90 FR 60970) · NYS Comptroller Report 2024-N-6, Dec 2025

The one real federal instrument that does exist measures something adjacent: the Census Bureau's biweekly survey of business AI adoption, which this project pulled directly and committed as a pipeline. It shows use climbing from 17.3% to 21.5% over eight months, with a three-fold spread across states. That is a real number, and it is not the number the policy debate needs. Knowing that a fifth of businesses use AI somewhere tells you nothing about whether an algorithm screened out your job application.

Census Business Trends and Outlook Survey, AI Use Supplement · pull_btos_ai.py, committed pipeline · 18 biweekly waves

PLAIN TALK

Nobody is counting. When someone tells you AI is causing mass discrimination in hiring — or tells you it definitely isn't — neither of them can prove it, because the United States does not collect the data that would settle it. The argument is loud, and the scoreboard is blank.

PART 2

The harms that actually reach a court are the boring ones

\$893M

LOSSES, FRAUD COMPLAINTS REFERENCING AI, FBI IC3, 2025

400K+

AI-LINKED CHILD-EXPLOITATION REPORTS, NCMEC, 2025

≈0

ADJUDICATED US CATASTROPHIC-AI CASES

~20–25

AI TORT CASES NATIONALLY

Against the measurement blackout, the adjudicated record is small but legible — and it is not where the discourse looks. Discrimination cases have produced actual settlements (the EEOC's first AI hiring settlement; a \$2.3M tenant-screening settlement). Fraud and synthetic child sexual abuse material dominate by sheer volume. The novel, catastrophic scenarios that organize most AI-policy attention have produced essentially no adjudicated American cases at all. Two different metrics, pointing the same way: whatever is generating legal consequences today is mundane.

FBI IC3 2025 Annual Report · NCMEC 2025 CyberTipline data · EEOC v. iTutorGroup · SafeRent settlement · ws03-findings

The obvious fix — let courts sort it out under existing law — is cheaper than a new regulatory regime, and it half-works. The evidence shows two distinct obstacles, not one. The first is access: in one 2026 case a magistrate judge refused to compel a vendor's AI bias-testing data as privileged; in another, a judge ordered broad discovery into a health insurer's coverage algorithm over its objection. Access is genuinely contested terrain, with real wins on both sides. The second obstacle is doctrinal and discovery reform would not touch it: product liability asks whether a product was defective *as designed and sold*, and a model that is retrained every few weeks has no fixed design to interrogate. And the case law is not accumulating the way a liability-first strategy needs — the leading generative-AI harm case settled in January 2026 before any ruling, which is a pattern, not an accident.

Mobley v. Workday (N.D. Cal.) · Estate of Lokken v. UnitedHealth (D. Minn.) · Garcia v. Character Technologies (M.D. Fla., settled Jan 2026) · ws03-findings

PLAIN TALK

The AI harms winning in court are discrimination and fraud — not robot apocalypse. And the "just let people sue" answer runs into two walls: companies can often keep the model itself out of evidence, and the law's basic question — "was this product defective when sold?" — doesn't fit software that rewrites itself monthly.

Nobody can check the safety homework

TRUST-BASED REGIME CONFIRMED

0

MECHANISMS WITH STANDING AUDIT ACCESS

C+

BEST LAB GRADE, INDEPENDENT 2026 INDEX

~20–30

STAFF AT THE FEDERAL AI EVALUATION BODY

Every major lab publishes a safety framework. Every one of those frameworks is self-assessed. The US AI Safety Institute was renamed in June 2025 to the Center for AI Standards and Innovation, its charter narrowed from broad safety to "demonstrable risks," and its evaluations remain voluntary — labs opt in. Independent technical evaluators do real work, but their access *in the US* is negotiated per engagement: API keys and compute credits, never standing rights to weights, training data, or internal deployment logs. The compute-reporting mandate that existed under the 2023 executive order was rescinded in 2025 and not replaced. No case was found, anywhere in the record, of an outside body catching a lab's published safety claim as false — not because labs are necessarily lying, but because no American body is positioned to check.

ws06-findings · CAISI charter and voluntary testing agreements · Anthropic RSP v3.0, OpenAI Preparedness Framework v2, Google DeepMind FSF v3.0 · Future of Life Institute AI Safety Index, Summer 2026

Correction, and it is the largest of this pass. This filing published that sentence with no country attached, and a Phase 2 steelman refuted the unqualified version out of the statute book. The EU AI Act's Article 92 lets the AI Office evaluate a general-purpose model, lets the Commission "appoint independent experts to carry out evaluations on its behalf," and lets it "request access to the general-purpose AI model concerned through APIs or further appropriate technical means and tools, *including source code*" — with fines of 3% of worldwide turnover, or €15 million, for refusing. Article 91 and Annex XI reach the other two items: training-data provenance, and "the computational resources used to train the model (e.g. number of floating point operations)." It binds the American labs named above, because they sell general-purpose models into Europe. Those powers became applicable on 2 August 2026. This filing was filed on 4 August 2026. **What is still true:** no exercise of that power is documented anywhere this pass could reach — no evaluation, no access request, no fine — and the provision names source code, not model weights. And in the United States nothing of the kind exists: California's SB 53 and New York's RAISE Act, the two frontier statutes actually in force, were read in full at primary tier and neither contains the word "audit" or grants access to anything.

Regulation (EU) 2024/1689, Arts. 51–55, 88–94, 101, 113 and Annex XI · Regulation (EU) 2026/1744 (Digital Omnibus on AI), which amends none of them · CA SB 53 as chaptered · NY RAISE Act, Ch. 699 of the Laws of 2025 · all committed in method/sources/ · steelman-log

The most useful evidence here is adversarial to the whole regime and comes from outside it: an independent index grades every major lab C+ or below and concludes — of Google DeepMind, OpenAI and xAI specifically — that stated commitments are an unreliable proxy for actual safety practice. That is not this desk's opinion — it is an outside body reaching the same conclusion this filing's hypothesis predicted. On the catastrophic-risk question itself, this filing declines to issue a verdict, and says why: expert estimates span from near-zero to near-certainty with no standardized elicitation method, and the surveys disagree partly because they are different instruments, not only because experts disagree. What the filing publishes instead is a ledger of what would change the assessment.

ws06-findings, "what would change this assessment" ledger · AI Impacts 2023 survey · 2026 Delphi catastrophic-risk study (different definition, not comparable)

PLAIN TALK

The AI companies grade their own safety homework. There is no inspector, no audit right, no way for anyone outside the building to confirm a lab is doing what it says. That's worth knowing whether you think the risks are overblown or underblown — right now, neither side can check.

PART 4

Your power bill is AI policy, and nobody files it that way

\$50–60B

RATEPAYER EXPOSURE CERTIFIED IN GEORGIA ALONE

\$29.4B

DATA-CENTER CAPACITY CHARGES, 4 PJM AUCTIONS

~\$10/mo

ATTRIBUTED TO DATA CENTERS ON D.C. BILLS

\$3–8B

TOTAL FEDERAL AI SPENDING (OBLIGATED)

The most consequential measurable AI-policy question in America right now is being decided in state utility proceedings, under a name nobody recognizes as AI policy. Georgia's commission certified nearly 9,900 MW of new generation with roughly 80% serving AI data centers, at an estimated \$50–60 billion in ratepayer cost over the assets' life — then froze base rates through 2028, an implicit admission the cost-shift risk was real enough to need a countermeasure. In the mid-Atlantic grid, PJM's own independent market monitor attributes \$29.4 billion across four capacity auctions to data-center load, including \$6.2 billion tied to data centers that have not been built.

Georgia PSC certification order, Dec 2025 · Monitoring Analytics (PJM independent market monitor), 2026 State of the Market · D.C. Office of the People's Counsel · ws05-findings

Utilities are building the fix: large-load tariffs in Georgia, Ohio, and Virginia now require data centers to commit to minimum charges of up to 85% of contracted capacity, with multi-year terms and exit fees. These are real cost-causer-pays mechanisms. They are also prospective only — they do not reach the generation already built and already in the rate base — and none has been tested against a large data center going bankrupt or never materializing. The comparison this filing had to assemble itself, because no source makes it: ratepayer exposure in Georgia and the PJM region alone (\$79–89B combined) far exceeds federal AI spending on the obligated reading (\$3–8B) and approaches it on the contract-ceiling reading (\$91.8B).

Georgia Power large-load rules · AEP Ohio tariff (PUCO, Jul 2025) · Virginia GS-5 rate class (eff. Jan 2027) · comparison assembled from separately-sourced figures, ws05-findings

PLAIN TALK

The biggest bill the public is actually paying for AI shows up on the electric statement, not in the federal budget. Georgia's regulators signed off on \$50–60 billion of it. It gets argued about as an energy story, almost never as an AI story.

PART 5

The jobs numbers have two parents

FORECASTS QUARANTINED

2

METHODOLOGICAL LINEAGES BEHIND MOST "% OF JOBS" FIGURES

13–16%

EMPLOYMENT DECLINE, EARLY-CAREER WORKERS IN EXPOSED ROLES

~0

AGGREGATE AI-ATTRIBUTABLE BREAK IN NATIONAL EMPLOYMENT

Trace the circulating "X% of jobs will be automated" numbers to their roots and they collapse into two families: the 2013 Frey-Osborne occupation-scoring method, now widely criticized and largely superseded, and the 2023 OpenAI-authored task-exposure method that underlies or partly underlies the IMF's headline figure and the major consultancies' frameworks. One genuinely independent instrument stands apart — the World Economic Forum's employer survey, a different kind of measurement entirely. Citing two of these as corroboration is citing one source twice. This filing therefore scores labor architectures against measured displacement only, with forecasts quarantined.

What the measurement shows is narrower and sharper than either the alarm or the dismissal. National employment shows no AI-attributable break; a New York Fed study found exposed and unexposed occupations already diverging *before* ChatGPT's release, with no additional break after. In freelance marketplaces, two independent research teams find real declines in writing and translation work. And one finding does not fit the calm aggregate picture: early-career workers aged 22–25 in the most AI-exposed occupations show a 13–16% relative employment decline in payroll microdata (16% in the paper's current version, 13% in its first), with software developers in that age band down about 20% from their late-2022 peak. That is narrow, real, and precisely targetable — and no architecture in the standard debate targets it.

NY Fed Liberty Street Economics, May 2026 · Hui/Reshef/Zhou, Organization Science 2024 · Teutloff et al., JEB0 2025 · Brynjolfsson/Chandar/Chen (ADP payroll microdata), 2025

PLAIN TALK

Most of the scary jobs numbers trace back to two studies wearing different hats. The real measured damage so far is small and specific: freelance writers and translators, and young people trying to start careers in tech. That last group is a real signal — and nobody's policy is aimed at them.

PART 6

Washington blocks; it doesn't build

99–1

SENATE VOTE STRIPPING THE STATE-AI MORATORIUM

2

FAILED FEDERAL PREEMPTION ATTEMPTS

0

FEDERAL STATUTES BINDING AI DEVELOPERS

Congress has twice tried to stop states from regulating AI and twice failed — a ten-year moratorium stripped from the 2025 reconciliation bill by a 99–1 Senate vote, and a similar rider dropped from the defense authorization after eighty-one House members objected. Having lost legislatively, the executive branch shifted to litigation: a December 2025 order created a Justice Department AI Litigation Task Force, which in April 2026 intervened against Colorado's algorithmic-discrimination law — the first time the federal government has gone to court to invalidate a state AI statute. Colorado's enforcement is suspended pending that case.

Senate Commerce record, Jul 2025 · EO 14365 (Dec 2025) · X.AI LLC v. Weiser, No. 1:26-cv-01515 (D. Colo.) · ws07-findings

One correction this filing owes its own first draft: the DOJ's actual theory is narrower than "the federal government opposes state AI regulation." Its complaint presses Equal Protection claims only — not the First Amendment, dormant Commerce Clause or vagueness theories the private plaintiff is arguing, and neither party pleads preemption at all. The remedy it seeks is not narrow, though: DOJ pleads the statute's duties inseverable and asks the court to declare all of SB24-205 invalid and enjoin its enforcement. And California's frontier-AI transparency law, Texas's TRAIGA, and Utah's amended act are all in force with no federal action against them. Colorado is so far the exception, not the pattern. What holds is the shape: federal activity to date has been aimed at stopping state law, not at replacing it with a federal rule that binds anyone.

DOJ Complaint in Intervention, ECF 17, *X. AI LLC v. Weiser*, No. 1:26-cv-01515 (D. Colo.) — read at primary tier in the 2026-08-10 verification pass · CA SB 53, TX HB 149, UT SB 149 as amended · ws07-findings

PLAIN TALK

There is no federal AI law. What there is: two failed attempts to ban states from writing their own, and a Justice Department now suing one state that did. Washington's main contribution so far has been to stop other people from acting.

PART 7

The coalition math: which fixes can actually pass

Four camps organize AI politics — safety organizations, acceleration advocates, civil-rights groups, and labor — and the standard assumption is that they deadlock everything. They don't, quite. Two architectures draw support from three of the four: protecting state-law primacy, and disclosure and transparency mandates. Only the acceleration camp sits outside both coalitions — opposed on state-law primacy, mixed on disclosure. But the safety camp's signature instrument, licensing frontier labs against compute thresholds, is genuinely isolated — the civil-rights and labor organizations have no located position on it at all, which is harder to build a coalition from than opposition would be.

ws08-findings coalition map: FLI/Encode, a16z/CCIA/Chamber, CDT/ACLU/AI Now, AFL-CIO · positions taken from testimony, public letters, official policy pages

One structural fact belongs in the open rather than in a footnote. The frontier labs fund a material share of the ecosystem that studies and evaluates them: an industry-funded safety fund makes grants to researchers who evaluate the funders' own models; the largest US AI-policy research center was built on a founding philanthropic grant and has taken industry money since; one lab runs a \$200 million fund for external research on AI's labor effects — one of this filing's own subjects. This is recorded as a disclosure column, not an accusation of capture. One evaluator's practice is worth naming as the contrast: it accepts model access but no cash from the labs it assesses.

Frontier Model Forum AI Safety Fund · Georgetown CSET founding and Google.org grants · Anthropic Economic Futures Research Fund · METR self-disclosed compensation policy

Two AI fixes could actually pass — letting states keep regulating, and requiring disclosure. The one the safety movement wants most, licensing the big labs, has almost no coalition behind it. Also worth knowing: the companies fund a good deal of the research about whether they're safe.

PART 8

What the old regimes teach — including the part that hurts

TRANSFERABILITY, HONESTLY

Four precedents, each with its disanalogy stated as plainly as its analogy. FDA premarket review is the model most AI-licensing proposals borrow, and its risk-tiering does transfer — but it evaluates a frozen product for a stated purpose, at roughly \$118 million per high-risk device and \$2.6 billion per drug. NEPA is the cautionary case: environmental review averages 4.5 years and has become a delay point somewhat independent of environmental outcomes, though its enforcement hook binds federal agencies, not private firms, and an AI disclosure law would not inherit it. China's algorithm registry is a real existence proof that a state can compel disclosure of training-data provenance — with about 45% actual coverage, only successes published, and enforcement leverage no US constitution supplies.

ws09-findings · Makower et al. medtech cost survey · Tufts CSDD 2016 · CEQ EIS timelines · CAC filing counts vs. independent model census

Nuclear regulation is the one that cuts both ways, and this filing owed it an honest hearing because its own hypotheses leaned skeptical of governability. The Nuclear Regulatory Commission's oversight process — resident inspectors on site, continuous inspection against a published manual, mandatory event reporting — is a working existence proof that continuous, independently verified oversight of a high-consequence technology is achievable. It is also bound up with the cost ratchet that helped end American nuclear construction: requirements tightened mid-build and never loosened. The two halves share a mechanism. You cannot import the verification without confronting the ratchet, and AI iterates in weeks where reactors took a decade.

NRC Reactor Oversight Process · Lovering et al., Energy Policy 2016 · older regulatory-ratchet cost literature (secondary compilation, provenance stated in ws09-findings)

A pattern surfaced across three of these independently, and it is the filing's quietest real finding: FDA approval, product-liability doctrine, and environmental review all assume a fixed artifact assessed at a moment. Continuously retrained software breaks that assumption in the same way each time. That is not three unrelated doctrinal quirks — it is one structural mismatch that AI governance keeps hitting from different directions.

Every regulatory model we might copy was built for things that hold still — a pill, a bridge, a reactor. The thing being regulated now changes every few weeks. Nuclear proves real inspection is possible; it also proves how that kind of oversight can strangle what it inspects.

PART 9

The scorecard, corrected

ARCHITECTURE	HARM	INNOV.	CATASTROPHIC	LABOR	CONCENTRATION
STATE-LAW PRIMACY + FEDERAL FLOOR	4	3	3	3	3
DISCLOSURE / TRANSPARENCY MANDATES	4	3	3	3	3
TARGETED HARM-CLASS STATUTES	3	3	3	3	3
LIABILITY-FIRST (DISCOVERY + DOCTRINE)	3	3	3	3	3
PUBLIC COMPUTE / OPEN MODELS	3	3	3	3	3
COMPREHENSIVE FEDERAL STATUTE	3	3	3	3	3
DO-NOTHING (COMPARATOR)	2	4	2	3	3
VOLUNTARY STATUS QUO	2	4	1	3	2
FRONTIER-LAB LICENSURE	2	2	3	3	2

Nine architectures, five objectives, six weightings. The red team earned its keep: the draft's rankings table did not match its own raw scores in a single one of the six columns, and two cells were scoring the wrong kind of evidence — one credited a proposed statute for the size of the problem it would target rather than any demonstrated effect, the other quietly gave the do-nothing comparator a better harm score than the status quo its own text called identical. All of it is in the log, with the corrected table hand-verified twice.

9 architectures × 5 anchored objectives × 6 weightings · adversarial red-team pass + independent blind re-score (11 cells) · ws10-scorecard, ws10-red-team-log, ws10-rescore-log

Corrected twice — a red team, then an independent blind re-score that moved eleven more cells — state-law primacy and disclosure/transparency mandates now co-lead most weightings, state-law primacy carried by the only large, multiply-sourced harm-reduction evidence in the entire

record, Illinois's biometric privacy law, which produced a \$650 million settlement and a company-wide facial-recognition shutdown where no federal guidance has moved a firm at all. Frontier-lab licensure shares the floor with the voluntary status quo on four of the six weightings — though not on the catastrophic-risk weighting licensure exists to serve, where the voluntary status quo alone holds the bottom; the comprehensive federal statute lands mid-pack as an all-neutral row, because the implemented EU evidence was read as supporting nothing above or below neutral on any axis — a reading the correction in Part 3 now contradicts on the statute. Those placements need their reason stated or they will be misread: they reflect scoring proposals *as actually drafted*, and no US frontier instrument — enacted or pending — requires anyone outside the developer to evaluate the model. Nuclear regulation proves a verifiable version is possible. No American bill is that version. The score is measuring the gap in what has been proposed, not a verdict that the goal is wrong.

Three of the four results just stated turn on a single scorecard cell each. A Phase 2 steelman contested four cells and published both readings rather than re-scoring on its own authority; the arithmetic is in a committed script that reproduces all fifty-four published values before computing a variant. Move **frontier licensure's concentration cell** from 2 to 3 — the correction the blind re-score already applied to the identical unmeasured claim in the comprehensive-statute row, and never applied here — and "shares the floor on four of the six weightings" becomes true on *none*. Move **disclosure's harm cell** from 4 to 3, the same correction the red team applied when it stopped a row being credited for the size of the problem it would address, and the co-lead ends: state-law primacy leads alone on all six. Move **the comprehensive statute's catastrophic-risk cell** from 3 to 4 on the EU evaluation power in Part 3's correction, and the row this page calls mid-pack co-leads four weightings and takes the catastrophic-risk weighting outright at 3.4. One result held under every variant tested: the voluntary status quo still scores worse than doing nothing, on all six weightings, even after the quotation behind its lowest cell was corrected to the narrower scope its source actually states.

ws10-scorecard sensitivity table · ws10/rank.py (committed; self-checks against the published v3 board) · steelman-log

CORRECTED VIA RED TEAM + BLIND RE-SCORE + PHASE 2 STEELMAN

APPENDIX

The honesty box

This is a Special Edition, and the label is doing real work. AI moves faster than a research cycle, so this filing stamps findings CONFIRMED, CONTESTED, or UNKNOWABLE-AS-OF-FILING and uses the third without embarrassment. **We do not adjudicate the catastrophic-risk question.** Expert estimates span orders of magnitude on instruments too different to compare;

the filing publishes a ledger of what would change the assessment instead of a probability. **The scorecard's first draft had broken arithmetic.** Its rankings table matched none of its six columns to its own inputs — caught by the red team, rebuilt, and hand-verified a second time before acceptance. **The AI-adoption baseline is single-root.** The Census survey is the only national instrument found; no independent second measurement corroborates its level. **Several primary sources were unreachable at filing — that disclosure has since been narrowed.** This box originally named the DOJ complaint, the federal courts' docket system, and a national lab's data-center load review as all refusing automated access. Two thirds of that no longer holds: the 2026-08-10 pass downloaded the entire *X. AI LLC v. Weiser* docket from RECAP without authentication and read the DOJ complaint (ECF 17) at primary tier — and two of the three claims that disclosure was hedging turned out to need correction. Still genuinely unobtained: the national lab's data-center load review and the Commerce Department's CAISI announcement, both of which still refuse automated access. Those claims, and only those, rest on convergent secondary reporting and say so. **One source in this domain did not survive checking:** a circulating account of EU AI Act "enforcement precedents" named no case and cited nothing; it was excluded. Assume there are more. **An independent blind second scorer has now re-scored every cell (2026-08-06)** — and corrected eleven of them, including one where this scorecard had dropped a cell below its own band definition, and one where the voluntary status quo earned a documented 1 that separates it, unfavorably, from doing nothing. The full reconciliation is in the record. **This filing has since been fact-checked against primary law (2026-08-10).** An independent pass under the verification protocol extracted 380 claims from these pages and the whole research record, recorded 380 verdicts, and corrected ten of them. The largest: this page published **"1.5M+" AI-linked child-exploitation reports for 2025 when NCMEC's own data says more than 400,000** — overstated roughly fourfold, now corrected here and in four record files. Also corrected: the healthcare disclosure rule is not a public registry; "over 96% of hospitals" is 96%, on 2021 survey data; the Justice Department's Colorado complaint asks the court to strike the *whole* statute, not one drafting choice; and frontier licensure does not share the floor on the catastrophic-risk weighting. What the pass confirmed at primary tier and did *not* change: the 99–1 Senate vote, the CFPB's 67 withdrawn guidance documents, the Colorado enforcement suspension, the Workday privilege ruling, the IC3 loss total, the comptroller's misrouted-calls finding, every Census figure, and all thirty-six values in the scorecard's rankings table. The three headline energy figures — Georgia's \$50–60B, PJM's \$29.4B, the D.C. bill impact — could **not** be reached at primary tier and remain on secondary reporting. Full ledger in the record. **A Phase 2 steelman then ran against the corrected filing (2026-08-10), in a separate session.** Its target was the position this filing's own protocol declared it owed a steelman to and then paid in a parenthetical: the case for governing frontier capability ahead of measured harm. It **partially won.** The largest thing it took: the sentence "no mechanism grants an outside party standing access to a lab's weights, training data, or compute logs" is withdrawn in its unqualified form and re-scoped to the United States, because the EU AI Act gives the Commission a compulsory, fine-backed power to obtain a general-purpose model and have independent experts evaluate it — applicable two days before this filing published, and never consulted. Also corrected: New York's RAISE Act is *enacted law*, not a pending proposal, and contains no third-party audit requirement — which

makes the American half of this filing's safety finding *stronger*. And four scorecard cells are now marked contested with both readings published, because three of the four scorecard claims on this page are one-cell results. What the steelman **failed** to take: the kill condition, the US safety finding, and the voluntary-versus-doing-nothing separation, which survives on all six weightings even at its own worst reading. A structurally blinded re-score of the four contested cells is owed and is recorded as owed.

THE RECEIPTS · 10-workstream protocol · Phase 0 gate with a GO verdict and a fired kill condition · pre-registered adjudication criteria per workstream · anchor table preserving broken priors (anchors 1 and 2) · deviations log (19 entries) · committed data pipeline (Census BTOS AI supplement, 18 waves, four cross-sections) · adversarial red-team pass with 11 logged attacks · independent blind re-score correcting eleven cells · Phase 1 fact-check against primary law (380 claims) · Phase 2 steelman with a committed ranking script · all public.

CITE · Gubment. "AI: We're Regulating a Quantity Nobody Measures." Special Edition Policy Whitepaper No. SE1, Aug 2026. gubment.com/ai · Sources & data: gubment.com/ai/sources