

To: Senator Josh Hawley (R-MO)

Policy Memo: The Making AI Governable for Americans Act (MAGA Act) [Apr. 2026]

Policy Recommendation: Congress should enact the Making AI Governable for Americans (MAGA) Act, requiring frontier AI developers, defined as those training models above 10^{24} floating point operations (FLOPs) or generating more than \$1 billion in annual AI-derived revenue, to (1) pay mandatory premiums into a federally administered AI Displacement Fund covering AI-motivated labor displacement, and (2) obtain pre-deployment security certification from a new, statutorily independent AI Security Evaluation Authority (AISEA) before commercial release.

Background: The core problem is moral hazard: frontier AI developers capture enormous private profits while offloading risk onto American workers and taxpayers. Goldman Sachs estimates AI could displace 6–7% of the U.S. workforce long-term [1, 13]. Early-career workers in AI-exposed occupations have already experienced a 13% relative employment decline since 2022 [2, 8]. The UK AI Safety Institute documents that frontier models can execute expert-level cyber attacks and accelerate bioweapon development, capabilities with direct national security and biosecurity implications [3, 4, 5]. Tort law partially addresses individual product liability, but diffuse harms like workforce displacement are structurally beyond its reach. State regulations, including California's SB-53 and New York's RAISE Act, have produced a fragmented patchwork that burdens businesses without closing the underlying accountability gap [6]. The Trump administration's 2026 National AI Legislative Framework preempts these state efforts while proposing no equivalent federal mechanism, leaving Big Tech to self-regulate on both security and displacement [7]. A handful of corporations make decisions affecting millions of American workers and families, with no accountability to either.

Policy Proposal: The MAGA Act addresses security and unemployment harms with a shared financial architecture: developer premiums fund both worker compensation and the regulator evaluating those same developers. Splitting the bill would double compliance overhead on an identical set of firms while severing the funding stream that keeps AISEA insulated from appropriations politics. Jurisdiction is triggered by either condition: models trained above 10^{24} FLOPs, or companies generating over \$1 billion in annual AI-derived revenue.

1) Mandatory Labor Displacement Insurance. Covered developers pay risk-adjusted premiums into a federally administered AI Displacement Fund. Premiums have two components: a *Base Assessment* set at a statutory floor of 1% of AI-derived revenue above \$1 billion, which Treasury may adjust upward to maintain a target reserve ratio, and an *Experience Rating Multiplier* applied firm-by-firm based on verified AI-attributable layoffs, mirroring unemployment insurance experience rating [9]. Amended WARN Act notices must designate whether AI systems substantially perform prior worker functions, with the substantiality

threshold set by Department of Labor rulemaking, enabling the actuarial feedback loop. Fund revenue supplements the Unemployment Insurance Trust Fund without new federal spending.

2) Mandatory Pre-Deployment Security Certification. No insurance payout remedies an extinction level event. Covered developers must obtain certification from AISEA before commercial release. AISEA maintains permanent technical evaluators supplemented by contracted red-teams for adversarial capability assessments, funded by a statutorily designated share of developer premiums rather than annual appropriations. The Trump administration's dismantling of the AI Safety Institute's mandate demonstrates why AISEA must be congressionally established with binding authority, insulated from executive restructuring [10, 11]. Board members serve staggered six-year terms and must hold demonstrated AI security expertise, ensuring no single administration can replace the full board. AISEA publishes certification criteria through notice-and-comment rulemaking, and denials are subject to Administrative Procedure Act judicial review, ensuring objectivity and recourse. Developers failing certification, domestic or foreign, are barred from U.S. commercial deployment, with U.S. market scale providing ample leverage, as GDPR demonstrated [12].

Implementation: The Department of Labor administers WARN amendments; Treasury administers the Fund and conducts annual actuarial recalibration. AISEA becomes operational within 18 months, with a 12-month premium phase-in. Certifications target a 90-day review window, with expedited 30-day review for incremental updates and a 30-day extension permitted for novel capability profiles. Post-deployment security incidents trigger automatic premium increases. To prevent cliff effects at the jurisdictional threshold, the base assessment applies only to AI-derived revenue above \$1 billion, mirroring progressive marginal income taxation. To prevent gaming, the Act applies to cumulative compute across model versions and to capability benchmarks set by AISEA rulemaking. AISEA conducts annual threshold reviews as training efficiency evolves, and below-threshold releases remain subject to mandatory capability disclosure.

Tradeoffs and Projected Impacts: Compliance costs fall disproportionately on large developers by design. There is a real risk that premiums become a cost of doing business rather than a behavioral deterrent. The Act mitigates this through the experience rating multiplier and Treasury's annual solvency-based recalibration authority, ensuring premiums escalate with realized harm; even if pricing fails to alter behavior, the Fund still shifts displacement costs from taxpayers to the corporations generating them. Mandatory certification may incentivize offshoring, but this concern is bounded: the U.S. market is too large for major developers to abandon, and the Act's jurisdiction applies at deployment rather than training location, so foreign-trained models still require certification. The alternative, no federal standard, leaves a fragmented state patchwork creating greater compliance uncertainty, while the shared funding architecture ensures both preventive and compensatory mechanisms adapt as capabilities evolve.

Citations

- [1] Goldman Sachs Research. (August 2025). “How Will AI Affect the Global Workforce?” goldmansachs.com.
- [2] Brynjolfsson, E., Chandar, B. & Chen, R. (August 2025). “Canaries in the Coal Mine? Six Facts about the Recent Employment Effects of Artificial Intelligence.” Stanford Digital Economy Lab Working Paper. digitaleconomy.stanford.edu.
- [3] Future of Life Institute. (Summer 2025). “2025 AI Safety Index.” futureoflife.org.
- [4] AISI / METR. (December 2025). “Frontier AI Trends Report.” aisi.gov.uk.
- [5] CSIS. (August 2025). “Opportunities to Strengthen U.S. Biosecurity from AI-Enabled Bioterrorism.” csis.org.
- [6] Orrick AI Law Center. (2025). “U.S. AI Law Tracker.” ai-law-center.orrick.com.
- [7] White House. (March 2026). “National Policy Framework for Artificial Intelligence: Legislative Recommendations.” whitehouse.gov.
- [8] Brookings Institution. Manning, S. & Aguirre, T. (February 2026). “Measuring U.S. Workers’ Capacity to Adapt to AI-Driven Job Displacement.” brookings.edu.
- [9] U.S. Department of Labor. (2024). “Trade Adjustment Assistance Program Overview.” dol.gov.
- [10] Brookings Institution. (March 2025). “A Technical AI Government Agency Plays a Vital Role.” brookings.edu.
- [11] TechPolicy.Press. (July 2025). “Renaming the US AI Safety Institute Is About Priorities, Not Semantics.” techpolicy.press.
- [12] European Parliament. (2016). “General Data Protection Regulation (GDPR).” EUR-Lex.
- [13] Acemoglu, D. (2024). “The Simple Macroeconomics of AI.” NBER Working Paper No. 32487.