

October 27, 2025

via eRulemaking Portal

(<https://www.regulations.gov/document/OSTP-TECH-2025-0067-0001>)

Ms. Stacy Murphy, Deputy Chief Operations Officer/ Security Officer

Ms. Ashley Lin, Strategic Coordinator

Office of Science and Technology Policy (OSTP)

The White House, Executive Office of the President (EOP)

National Coordination Office (NCO) 2415 Eisenhower Avenue, Alexandria, VA 22314

Re: Regulatory Reform on Artificial Intelligence - 2025-18737 (90 FR 46422)¹

Dear Ms. Murphy and Ms. Lin,

On behalf of Data Boiler Technologies, I am pleased to provide the OSTP with our comments on the captioned release to “Regulatory Reform on Artificial Intelligence (AI)”, thereafter, referred to as the “RRAI”. As an inventor of patented solutions (US, Canada, Singapore, Japan and recently approved in Australia and EPO) in signal processing, ensemble learning, trade analytics, time-lock cryptography, etc., we are among the first to pin point the true “*AI Risks*”² (= downfall of humanity, including but not limited to the possibility of AI taking down the energy grid). We suggest the OSTP to review this comment letter in conjunction with our submitted comments³ to the National Science Foundation (NSF) - the Networking and Information Technology Research and Development (NITRD) of National Coordination Office (NCO) in March 2025 on the Development of an AI Action Plan.

Key takeaways

Revise the definition of AI under 15 U.S.C. § 9401(3) to preempt state AI laws, suggest using the following or the likes:

“Covered AI technologies refer to cognitive systems (beyond learning from pairing a neutral stimulus that becomes a conditioned stimulus), comprise of memory AND topology of known lessons, that learn from regularities and irregularities of pattern(s) / knowns and unknowns / models/ simulations, AND

the system’s internal process EITHER comprises of multi-steps reasoning (understand in a way that mimics humans; NOT merely extracting signals to generate alerts; NOT unconscious thinking) OR capable of generating datum uniquely different from a plagiarized copy, that

manipulates or presents at least an abstracted phenomenon (person or avatar, thing or computer-generated element, or real or virtual event that is hypothetical or observed to exist or happen in a distanced past, real-time, or irrespective of spacetime) in a metaverse, real, or virtual environment

autonomously OR follow commands/ instructions, to generate expectations, make-believe, or assert that certain selected or perceived phenomenon is or will occur / available for use (regardless of the system internalizes, consumes, or makes feed(s) / datum available to its users in a domain, a dark web, or any iteration of the internet or intranet), AND

through action (including provision of customized or generic recommendation that reinforces, strengthen or weaken an ideology) OR inaction to stimulate the thought processes of at least an individual human OR the operations of a machine.”

¹ Docket ID#: OSTP-TECH-2025-0067 <https://www.federalregister.gov/documents/2025/09/26/2025-18737/notice-of-request-for-information-regulatory-reform-on-artificial-intelligence>

² <https://www.linkedin.com/pulse/ai-risks-downfall-humanity-kelvin-to-yhzge/>

³ https://www.databoiler.com/index_htm_files/DataBoiler%20NSF%20NITRD%2020250315.pdf



Followings are our answers to specific questions:

(i) What AI activities, innovations, or deployments are currently being inhibited, delayed, or otherwise constrained due to Federal statutes, regulations, or policies? Please describe the specific barrier and the AI capability or application that would be enabled if it was addressed. The barriers may directly hinder AI development or adoption, or indirectly hinder through incompatible policy frameworks.

Repeal of the last administration's executive orders 14110 on *"safe and trustworthy AI"* realigns the US AI policies from a [GMMM](#) risk management framework to a pro-growth priority. Bureaucrats temporarily lose direction when *"governance"* is no longer the center of everything. The so-called *"standards"* for algorithmic audits and bias mitigation derived from long list of OECD / ISO *"global AI requirements"* remain haunting the industry (see [point 8](#) and [point 11](#)). They do NOT fit the US best interest and are barriers affecting US AI firms' competitiveness in overseas.

We are skeptical and have reservations to certain high-risk AI projects (e.g. facial recognition, predictive policing, hiring algorithms) that are incompatible to Vice President JD Vance's remarks about *"American AI will not be co-opted into a tool for authoritarian censorship"* or they contradict the current administration's merit-based policy (EO 14173) that dismantles DEI initiatives. Do not get us wrong, we strongly support the use of AI to detect illicit activities for example, those involving Digital Assets / DeFi / De-dollarization. Instead of *"massive government surveillance"* that invades privacy, we have innovative methods to accomplish the regulatory goals. We recommend Federal preemption of STATE AI and privacy laws to reduce regulatory complexity. For that, policy makers should build on the commendable US copyright and AI report. It strikes the appropriate balance in that assessing the divergence between private rights and social costs.

Hindrance mainly comes from false premises where big laws/ consulting firms prior lobbying efforts in attempt to sell existing GRC, Business Continuity, Resiliency, Cybersecurity, Privacy, and non-discriminatory tools, and then regurgitate them as the *"Foundations of a responsible AI risk management framework"*. AI risk management approach that puts *"Govern"* in center of *"Map, Measure, and Manage"* is oversimplified. We recommend equating AI risks to the downfall of humanity to streamline the framework and relevant guidance.

Push back on foreign countries' AI Acts deem to be their protectionism policy. They create unnecessary bureaucracy and favor subjective judgements. The long lists of *"global AI requirements"* / *"standards"* are barriers affecting US AI firms' competitiveness in overseas. In order for the US to exert influence on Global AI policies, US AI regulatory regime must be smart and precise (contrasted to the newly released China's AI Safety Governance Framework 2.0).

We welcome RRAI removal of barriers, promote faster innovation, and ensure American dominance in the global AI race.

(ii) What specific Federal statutes, regulations, or policies present barriers to AI development, deployment, or adoption in your sector? Please identify the relevant rules and authority with specificity, including a cite to the Code of Federal Regulations (CFR) or the U.S. Code (U.S.C.) where applicable.

We applaud the US Securities and Exchange Commission (SEC) for formally withdrawing the problematic proposal from the last administration about *"Predictive Data Analytics"*⁴ that was detrimental to AI innovation. Per our 2023 submitted comments to the SEC,⁵ we argue that *"the scope is overly broad and would become a tollgate every time an investment firms procures new technology, they will have to first consult a law or consulting firm – result in, the 'non-TECH bureaucrats' regulating the 'TECH professionals' and corruption."*

⁴ <https://www.sec.gov/files/rules/proposed/2023/34-97990.pdf>

⁵ https://www.databoiler.com/index.htm_files/DataBoiler%20SEC%2020231010%20Predictive%20Analytics.pdf

(iii) Where existing policy frameworks are not appropriate for AI applications, what administrative tools (e.g., waivers, exemptions, experimental authorities) are available, but underutilized? Please identify the administrative tools with specificity, citing the CFR or U.S.C. where applicable.

The following 2 Acts are inappropriate:

- AI in Government Act of 2020: Requires federal agencies to develop governance frameworks before deploying AI. This has slowed adoption due to lengthy compliance reviews and interagency coordination requirements.
- Advancing American AI Act: Mandates risk assessments, transparency, and human oversight for federal AI systems. Agencies must meet 94 distinct requirements, creating friction for pilot programs and procurement.

The White House's AI Action Plan includes many good recommendations. However, most are still under review or pending interagency coordination. Government shut down also affects progress. Amid Federal agencies may be constrained by procurement rules that limit rapid adoption of AI tools, certain staffs have negative sentiment against AI given the Department of Government Efficiency (DOGE) has stormed their departments with AI that unveiled productivity concerns. Some only go to the large laws / consulting firms for AI recommendations instead of immerse in learning related TECH.

We favor usage of a sandbox approach to provide a safe environment for AI / FinTech vendors. Thus, allowing thoroughly testing of products while awaiting various authorities overhauling the regulatory regime on AI and oversight of BigTECH. As long as there is no cheating, lying, stealing, and/or violations of export regime, US Patriot Act, FinCEN rules, and CFAA, granting sandbox participants immunity or certain limited liability may help. Sandbox in turn would be informative to guide future direction of policies development.

To support local innovations, the US Patent and Trademark Office (USPTO) should take a more liberal approach. Amid the available of subject matter eligibility guidance and AI related updates⁶ explains how the USPTO personnel should evaluate claims for patent subject matter eligibility under 35 U.S.C. 101, patent examiners follow-through on doing their job objectively and diligently remain a problem or doubtful. The Memorandum⁷ by the USPTO Deputy Commissioner Kim seems weak or insufficient to correct subjective bureaucratic behaviors that hinders the development of AI in the US.

We hope Congress can quickly adopt our suggested re-definition of AI under 15 U.S.C. § 9401(3) to preempt state AI laws. Internationally, we are afraid the US being late in shaping the right AI Global Policies, letting countries like China get much attention to persuade their agenda (see [our response to Question iv](#) and [point 10](#)).

AI often involves huge investments, except copycats that free ride from "*knowledge distillation*". To support US-based AI firms to compete with foreign STATE-sponsored rivalries, we and other small AI developers would require diplomatic support before considering setting ourselves abroad in sensitive TECH such as AI.

Not ONLY does the US need better framework than CN-AISGF2, the US standards have to be uplifted by mass adoption around the world. When part of the world may not be abiding to the Bretton Woods - US lead formal orders, an effective way for the US to deter de-dollarization movements or other illicit finance activities, is the emergence of informal sub-orders, inclusive of the use of unconventional pressurization and other diplomatic approaches (e.g. *given FinCEN newer authorities, similar to Section 311, in Section 2313a of the Fentanyl Sanctions Act and Section 9714 of the Combating Russian Money Laundering Act*) to pragmatically push foreign nations to go back to the negotiation table and ensure American dominance in the global AI race.

⁶ https://www.uspto.gov/sites/default/files/documents/OCE-DH_AdjustingtoAlice.pdf

⁷ <https://www.uspto.gov/sites/default/files/documents/memo-101-20250804.pdf>



(iv) Where specific statutory or regulatory regimes are structurally incompatible with AI applications, what modifications would be necessary to enable lawful deployment while preserving regulatory objectives?

Aside from our suggested revision to the definition of AI under 15 U.S.C. § 9401(3) on [page 1](#) of this comment letter, we recommend a complete overhaul of US NIST-AIRMF 2022. Particularly, [GMMM](#) is oversimplified. Remove the requirement about “neutralize” biases that pursuit of consensus. “Bias” in competitive settings is essential. Fairness can dilute the US strategic advantage, especially when foreign adversaries are not playing by the same rules. It is a paradigm shift to go from suspicions to opportunistic about newfound signals among chaos, AI “hallucinations” may discover unknown unknowns which were previously nonsensical to human.

OECD’s definition of AI system seems overly focused on Large Language Models (LLMs) and undermines other cognitive reasoning AIs. One-size-fits-all global standards or lack sector tailoring is against big trends towards mass customization that unleash tremendous values. Standards setters manufacture increasing number of new certificates to sell their publications, training, assessment, audit, and consulting services. Despite the high degree of overlap, they are hardly incorporated into streamlined versions in practice. In turn, accumulated bureaucracies exacerbate the gap between the “Haves” and “Have-Nots”.

ISO 5338 positions AI as an extension of existing software lifecycles—not a standalone domain. It could be incorporated in ISO/IEC/IEEE 12207, where ISO/IEC/IEEE 24748-1 may serve as meta-framework that harmonizes software, systems and AI lifecycles. There are too many details that are applicable to all technologies. Rights to revoke a user agreement with standard provisions, such as “no illicit or manipulative use of technology” may suffice. The US does have CFAA amid its narrow scope does NOT impute liability to internal workers who disregard a use policy.

The US RRAI should lead the world to re-center the focus on our identified key AI risks (energy; addictive, herd and/or polarized behaviors / destroy humans’ abilities to think independently; censorship; hyper optimization; insurgent / unhealthy competition) to mitigate downfall of humanity. Be cautious on countries with human rights violations that hide under the guise of “safety monitoring and assessments”. Policy Makers should consider the [Asimov’s Three Laws](#) and [Zeroth \(Forth\) Law](#) for AI. The ISO 23894 Risk, ISO 42001 management system, and ISO 38507 governance frameworks ought to be redirect accordingly. Alter or get rid of ISO 24028 trustworthiness, ISO 24027 bias, ISO 5469 functional safety. Correct ISO 22989 terminology. Upon proper overhaul, more AI firms may be interested in getting certified.

How AI firms modularize their machine learning workflows, piping their data, engineer their reasoning algorithms and whatnot in ISO 23053 are private matters. The foundation of a responsible AI is NOT about how good one person can articulate or reveal the secret ingredients of an AI to others. In general, we do support interoperability. Innovations such as Model Context Protocol works like a USB-C. How it may be used together with Agent AI is a wait and see. Varying protocols have different pros and cons. Regulators may only step-in if Elites are holding essential datum within their close network and restricting public access that detriment to market efficiency or trigger Antitrust concerns.

EU’s GAPI is problematic. How they would consider “the Code’s design and build coherence where needed”? “Coherence” limits creativity. Differentiation is what drives innovations. US RRAI should heed the lesson and avoid policy mistakes.

(v) Where barriers arise from a lack of clarity or interpretive guidance on how existing rules cover AI activities, what forms of clarification (e.g., standards, guidance documents, interpretive rules) would be most effective?

Clarifies the Federal Trade Commission (FTC)’s 2022 guidance on “algorithmic unfairness” interpreted §5 of the FTC Act.⁸

⁸ <https://www.americanbar.org/content/dam/aba/publications/antitrust/magazine/2023/vol-37-issue-3/ftcs-2022-policy-statement-scope-unfair-methods.pdf> ; https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5212510

(vi) Are there barriers that arise from organizational factors that impact how Federal statutes, regulations, or policies are used or not used? How might Federal action appropriately address them?

The US Government Accountability Office (GAO) May 2025 report – *AI: use and oversight in Financial Services*⁹ reference an investor Education and Advocacy material by the SEC and FINRA,¹⁰ but completely miss the problematic proposal from the SEC last administration about *“Predictive Data Analytics”* that was detrimental to AI innovation (*the proposal was formally withdrawn on June 12*). The report merely suggests and reiterates a 2015 recommendation that pushing the ball to Congress to consider granting NCUA authority to examine technology service providers for credit unions and ask NCUA to update model risk management guidance. The disconnect and outdatedness seem ONLY the tip of the iceberg.

GAO September 2025 report,¹¹ highlighted that agencies operate under 94 AI-related requirements and face oversight from ten executive branch groups. Mentality of *“the more you do, the more prone to challenges in fulfilling these long-list of requirements; hence, tendency to do less or nothing in pursuing of AI”* may reflect bureaucracy runs deep within US Government, amid we do recognize some hardworking folks do exist at selected agencies.

Drastic reform is necessary to address organizational factors that impact how Federal statutes, regulations, or policies are used or not used. The problem appears bigger than the following:

- a. Fragmented leadership
- b. Risk-Averse Culture and Compliance Paralysis
- c. Procurement Bottlenecks
- d. Lack of Interagency Collaboration
- e. Underutilized experimental authorities
- f. Swing in policy directions

Calling for a mandate to appoint Chief AI Officers across all major agencies, and creating a centralized AI Coordination Council to align efforts across agencies seems insufficient to drive essential change. We recommend DOGE to step-in. Prioritize clear procurement hurdles to allow outsourcing of certain Government functions to be run by private sector at lower cost with better efficiency. Leverage AI and automated systems to replace an unfit workforce is inevitable.

Respect human’s creativity and ingenuity – Enhance Copyright Laws to aligns AI rights and obligations

We applaud the White House Special Advisor Mr. David Sacks’ speech on June 10.¹² Global proliferation of American AI technologies and President Trump’s pro-growth stance are encouraging. We agree with his observations *“we’re moving from models that are essentially chatbots to agents that can take actions on the user’s behalf ... liberation of these digital agents ... and those agents are connected ...”* We are glad *“the Federal Government will serve as enabler of Research and Development ... friendly regulatory policy and more data center and energy resources.”* The following pages showcase how we formulate our recommendations for meaningful scope of AI rules that address regulatory mismatches; reduces structural incompatibility; improves regulatory clarity; removes direct hindrance on AI development, deployment, and adoption; as well as supports organizational factors that will positively influence workforce readiness, institutional capacity, or cultural acceptance of AI adoption.

⁹ <https://www.gao.gov/assets/gao-25-107197.pdf>

¹⁰ <https://www.finra.org/investors/insights/artificial-intelligence-and-investment-fraud>

¹¹ <https://www.gao.gov/assets/gao-25-107933.pdf>

¹² https://youtu.be/XhdOA3_e8Nw?si=Jgjl4iluv8XHlY5O

1. AI is NO ordinary “*machine-based systems*” or “*automations*”, but a “*cognitive system*” capable of “*learning*” to continuously improve the Functioning of a Computer / to any Other Technology or Technical Field (e.g. GPS Satellite system is NOT AI, but the traffic prediction and personalized recommendations are).
2. By referring AI as “*cognitive*”, the focus is on the system’s **internal** (mental like) processes, **rather than SOLELY on external** behaviors and consequences. Learning per “*Dog Salivating Theory*”¹³ that nurtures voluntary behavior by pairing external conditioned stimulus to associate two or more phenomena, such technique alone should NOT be constitute as AI. Computational techniques that mimic pet training in itself is insufficient or unlikely to inflict harm on humanity. Also, the presence or absence of “*operant conditions*”¹⁴ could merely be automation in itself to modify voluntary behaviors through external consequences (reacting to the laws) of rewards and punishments. If without combining with other internal (mental like) computational techniques, such systems should NOT be constituted as AI.
3. CAPTCHA¹⁵ – a security test that uses a “*Turing test*”¹⁶ to differentiate between humans and bots, is NOT an AI itself. However, Google reCAPTCHA system is part of an AI that **consists of internal processes** (e.g. keylogging to track and analyze user interactions, such as mouse movement and typing patterns) to determine if a user is human. Keylogging¹⁷ if without users’ consent could be invasive to **privacy**, hence it should be a regulated activity to prevent it from inflicting harm on humanity. Another computing activity that should be regulated is – the use of AI to help bots bypass CAPTCHA or the like security test,¹⁸ except ethical hacking.
4. According to the Department of Justice (DOJ) with regard to the **Computer Fraud and Abuse Act** (CFAA)¹⁹, “*In either a ‘without authorization’ case or an ‘exceeds authorized access’ case, the attorney for the government must be prepared to prove that the defendant knowingly accessed a computer or area of a computer to which he was not allowed access in order to obtain or alter information stored there, and not merely that the defendant subsequently misused information or services that he was authorized to obtain from the computer at the time he obtained it.*” 18 USC §1030 can **curb** activities such as: using AI to automate **unauthorized access**; training AI on **data obtained via unauthorized scraping** of protected websites or systems; deploying AI bots (agents) that **exceed authorized access** via e.g. brute-force login attempts, **credential stuffing**; and “*prompt injection*”²⁰ (exploiting LLMs struggle to differentiate between instructions from the developer and malicious input from a user, where attacker override instructions via direct insertion of malicious commands or indirectly embed hidden malicious instructions to alter behavior of a system) attacks that **bypass content filters or access controls**.²¹ Separate AI law for the same crime seems unnecessary.
5. Penn State Law Professor Ido Kilovaty argues that bypassing a code-based restriction through “*prompt injection*” is a form of access without authorization.²² In Van Buren v. United States (2021),²³ “*using a system for an improper purpose (e.g., personal gain) does NOT violate CFAA if the user had legitimate access to the system.*” Unlike China’s laws on dishonest computer use that is part of their centralized state control over data and cyberspace against threat

¹³ <https://cavcanine.com/classical-conditioning/>

¹⁴ <https://psychcentral.com/health/operant-conditioning>

¹⁵ <https://en.wikipedia.org/wiki/CAPTCHA>

¹⁶ <https://www.techtarget.com/searchenterpriseai/definition/Turing-test>

¹⁷ https://en.wikipedia.org/wiki/Keystroke_logging

¹⁸ <https://www.sciencefocus.com/future-technology/ai-vs-captcha>

¹⁹ <https://www.justice.gov/jm/jm-9-48000-computer-fraud>

²⁰ <https://www.ibm.com/think/topics/prompt-injection>

²¹ <https://www.lawfaremedia.org/article/when-manipulating-ai-is-a-crime>

²² <https://www.lawfaremedia.org/article/when-manipulating-ai-is-a-crime>

²³ <https://natlawreview.com/article/scotus-resolves-circuit-split-limits-scope-computer-fraud-and-abuse-act>



to national security and social stability to severely punish so-called “*unauthorized activities*” subjectively, CFAA is a much narrower statute if compares to other jurisdictions.²⁴ CFAA meant to target external hackers’ unauthorized access and damage, it does NOT impute liability to internal workers who disregard a use policy.

Stanford Law Professor Orin Kerr in his essay²⁵ suggests that **bypassing code would constitute a cybercrime ONLY if the code is a “real barrier” as opposed to a “speed bump.”** With increasing algorithmic activities directly and indirectly interacting with another algorithmic activities, e.g. matching queues in electronic trading are constantly interfered by different tactics – be it using low latency advantage to get ahead, deliberate use of “*speed bump*” to slow things down, detection of adverse selection, avoidance of toxic, slippage, alternation of bandwidth, etc. – all being permissible under securities laws, what is and how would the “*Norm*” **evolves overtime** in contriving “*real barriers*” to defend against offensive “*penetration*” e.g. quote stuffing to gain edge over competitors (if NOT spoofing)? Bureaucracy accumulates if policies favor protectionism (e.g. whether order protection can be replaced by competing market forces,²⁶ the “Octopus” is hard to untangle). Rather than calling it a moot point, remember that there will **always be a technologies arm race between offense and defense.**

6. We agree with the NIST AI Risk Management Framework (NIST-AIRMF) Playbook (page 63 - Map 1.4)²⁷ where it said, “*Socio-technical AI risks emerge from the interplay between technical development decisions and **how a system is used, who operates it, and the social context into which it is deployed.***” Yet, that is applicable to all technologies, **not just AI** (e.g. airplanes to fulfill human’s dream of flying turn into jet fighters for war). We do encourage embedding appropriate controls in tech product designs. Whilst requiring AI firms to “*establish comprehensive and explicit enumeration of AI systems’ context of business use and expectations*” and asking examiners to understand every bit of “*contextual factors may interact with AI lifecycle actions*” is overkill. Rights to revoke a user agreement with standard provisions, such as “*no illicit or manipulative use of technology*” may suffice.

Original inventor may **never come up with an exhaustive list** of usage purposes or anticipate possible repurpose of his/her technology. Free enterprise should NOT be obligated to reveal its **secret ingredient** of their technologies, **unless being identified with evidence for suspicious crime.** How “organizational mission and identified system purpose create incentives within AI system design, development, and deployment tasks that may result in positive and negative impacts” may be irrelevant, if AI firms are not restricted to expand its **economy of scope** under the pro-growth AI policies. We criticize the FINRA Consolidated Audit Trail (CAT) system,²⁸ particularly “*the defined purpose of accessing CAT should be much narrower than the broadly defined ‘regulatory purpose’ ... there should be no access to CAT for ‘market surveillance’ purpose prior to identifying symptoms of irregularity that are substantiated by ...*” for our concerns about “*function creep*”²⁹ and the realism of various adverse scenarios.³⁰ Some degree of scope limitations based on **scale** of operations is recommended, amid this is applicable to all products and technologies, **not just AI.**

7. A crawler or web scraper to **automatically** extract data from **external** environment (web) is NOT inherently an AI. When combining crawler’s function with **internal** processes to for example, involving **intelligent data analysis** to **enhance** the data collection with **contextual understanding** (rather than just the keywords search), then such system is an AI. Another example – scanner or camera to surveil the public area outside of private property is NOT an AI itself.

²⁴ <https://digitalcommons.harrisburgu.edu/cgi/viewcontent.cgi?article=1006&context=other-works>

²⁵ <https://columbialawreview.org/wp-content/uploads/2016/05/Orin-S.-Kerr.pdf>

²⁶ <https://www.linkedin.com/pulse/can-order-protection-replaced-competing-market-forces-kelvin-to-qfgte>

²⁷ https://www.acc.com/sites/default/files/2024-08/2023.01.26-AI_RMF_Playbook.pdf

²⁸ https://www.databoiler.com/index.htm_files/DataBoiler%20SEC%20CAT%2020210503.pdf

²⁹ <https://www.lawinsider.com/dictionary/function-creep>

³⁰ https://en.wikipedia.org/wiki/Edward_Snowden ; https://en.wikipedia.org/wiki/2015%E2%80%932016_SWIFT_banking_hack

Sophisticated system that has internal processes to control one or orchestrate multiple surveillance camera(s) to enhance the monitoring with contextual awareness (e.g. facial recognition to analyze identity) is an AI.

8. We are thankful for Vice President JD Vance remarks³¹ at the AI Action Summit, in particular *“AI must remain free from ideological bias, and that American AI will not be co-opted into a tool for authoritarian censorship.”* It helps address civic concerns over massive government surveillance.³² NOTE: *“Ideological bias”* is a human bias driven by political or social belief. Ideology may persuade citizens to believe one ideology is better than the others. Foreign adversaries attempt to tarnish American ideology for their envy of the US advantages in achieving more with less. Their provocations fail because any ideology is the enemy of free will in the US long standing heritage.

Whereas *“bias”* in many domains – especially competitive ones like defense or finance – bias is not just inevitable, it is essential to: prioritize certain outcomes (e.g. speed over accuracy, stealth over transparency), reflect strategic performance (e.g. risk tolerance, adversary modeling), and/or to exploit asymmetries (e.g. alpha in trading, deception in war). There are different AI machine learning algorithms, some use cognitive reasoning for multi-steps strategic plan (e.g. chess game) where *“bias”* is essential, others use non-reasoning (or generative) models excel at fast, pattern-based tasks like content generation or chatbots where consensus building and/or optimization for the most commonly accepted respond (consistency in reproducibility of outcomes) is prioritized. One size does not fit all.

Unfortunately, the last administration assumed or interpreted *“AI bias”* as systematic and repeatable error in a computer system that creates unfair outcomes, such as disadvantaging a particular gender or race. This contradicts with the current administration’s merit-based policy (EO 14173)³³ that dismantles Diversity, Equity and Inclusion (DEI) initiatives. Trying to *“neutralize”* biases in pursuit of consensus or fairness can dilute the US strategic advantage, especially when foreign adversaries are not playing by the same rules.

9. Amid the EU AI Act is the world’s first AI law calling for *“shared governance”* in attempt to *“shape and inform effective Global governance of AI”*, we do NOT see that the world is following their footsteps. Their mandate of development of a Code of Practice on General-Purpose AI (GAPI) is problematic. Best practice sharing is not wrong. Concern is – how they would consider *“the Code’s design and build coherence where needed”*? Regurgitates GRC tools as AI compliance is the wrong approach. *“Coherence”* limits creativity. Differentiation is what drive innovations.

European Commission’s recent preliminary founding of TikTok and Meta in breach of their transparency obligations under their Digital Services Act (DSA)³⁴ may be a diplomatic move to *“test the water”* on whether China and/or US would *“cooperate”* and respect their authority. We do recognize that “Big-TECH” could make it cumbersome for anyone, including users’ denial to share their own profile and activities data with the service provider’s integrated partners, cancellation of services, as well as researchers to request access to public data. However, their requirements lack clarity. Online marketplaces, social networks, content-sharing platforms, app stores, and online travel and accommodation platforms are like news media, where acclaimed author Alain de Botton said, *“The news, while attempting to inform, often selectively highlights certain aspects rather than recording everything in its entirety.”*³⁵

Their Digital Markets Act (DMA) aims to ensure *“fair competition and practices among large online ‘gatekeeper’ platforms like search engines and app store”* is nothing but protectionism policy. Invoke Antitrust laws may suffice.

³¹ <https://www.presidency.ucsb.edu/documents/remarks-the-vice-president-the-artificial-intelligence-action-summit-paris-france>

³² <https://cs.stanford.edu/people/eroberts/cs181/projects/ethics-of-surveillance/ethics.html>

³³ <https://public-inspection.federalregister.gov/2025-02097.pdf>

³⁴ https://ec.europa.eu/commission/presscorner/detail/en/ip_25_2503

³⁵ <https://www.amazon.com/News-Users-Manual-Alain-Botton/dp/0307379124>

10. China recently released their “AI Safety Governance Framework 2.0” (CN-AISGF2).³⁶ “China’s golden age of hacking”³⁷ reflects their intention to rival against the US leadership in AI development. Their cybersecurity law,³⁸ computer crime criminal law (Articles 285-287),³⁹ and Personal Information Protection Law⁴⁰ are their broader policies that prioritize their national security and state control. It lacks coverage on the appropriate delineation of private rights and social costs.⁴¹ The table below is a quick comparison of China’s latest AI frameworks vs the US NIST-AIRMF 2022 version:

	CN-AISGF2	US NIST-AIRMF 2022
Governance philosophy and binding nature	Top-down, state-led with strong regulatory oversight. Non-binding but aligned with national policy priorities. Carnegie Endowment for International Peace expects China’s system is likely to translate it into binding technical standards and regulatory tools ⁴²	Voluntary bottom-up guidance for organizations. Sector-neutral – we are concerned it’s a “one-size-fits-all approach.”
Implementation tools and structure	Multi-tiered risk classification with corresponding safeguards and lifecycle controls. Encompass the full lifecycle from design to deployment, and monitoring. They suggest direct interventions at the model layer with risk classification by model type, call for documentation tracing the origin of training data to ensure adhere to China’s content safety and national security standards, measurements for frontier safety, and encourage embedding filters and constraints within the model’s architecture or decoding logic. “Strengthen assessment on Frontier safety, downstream propagation and amplification of model defects” reflects their paranoid / mistrust of US lead advance AIs, amid “knowledge distillation” ⁴³ (copycat) accelerated their AI development.	Putting “Govern” in center of “Map, Measure, and Manage” (GMMM) for every items. It looks more like an unfinish laundry list, awaiting additional works to organize and turn it into various guidances on tradeoffs and relationships that may exist among trustworthiness characteristics, explainability, interpretability, risk tolerances, tutorials, etc. Regurgitate GRC tools as AI compliance only benefits law / consulting firms. It does NOT help address true AI risks.
	CN-AISGF2 looks undeniably comprehensive. Their usage of familiar GRC best practices make it appealing for foreign jurisdictions to adopt. The US NIST-AIRMF playbook and related guidances if blindly continuing the oversimplified GMMM path may end-up similar to CN-AISGF2. In order for the US to exert influence on Global AI policies, US AI regulatory regime must be smart and precise. We recommend equating AI risks to the downfall of humanity. Center on our identified key risks, ² and build on this commendable <u>US copyright and AI report</u> . ⁴⁴ It resolves an increasing number of AI copyright lawsuits, ⁴⁵ aligns rights and obligations that respect human’s creativity and ingenuity, fits the simple but effective <u>Asimov’s Three Laws</u> ⁴⁶ and <u>Zeroth (Forth) Law</u> , ⁴⁷ and commands the world’s order to follow the US lead.	

³⁶ https://www.cac.gov.cn/2025-09/15/c_1759653448369123.htm

³⁷ <https://www.npr.org/2025/07/19/nx-s1-5471340/why-this-is-chinas-golden-age-of-hacking>

³⁸ <https://digichina.stanford.edu/work/translation-cybersecurity-law-of-the-peoples-republic-of-china-effective-june-1-2017/>

³⁹ <https://www.warnathgroup.com/wp-content/uploads/2015/03/China-Criminal-Code.pdf>

⁴⁰ <https://digichina.stanford.edu/work/translation-personal-information-protection-law-of-the-peoples-republic-of-china-effective-nov-1-2021/>

⁴¹ <https://iea.org.uk/wp-content/uploads/2016/07/THE MYTH OF SOCIAL COST.pdf>

⁴² <https://carnegieendowment.org/research/2025/10/how-china-views-ai-risks-and-what-to-do-about-them?lang=en>

⁴³ <https://www.sciencedirect.com/science/article/pii/S2666827024000811>

⁴⁴ <https://www.copyright.gov/ai/Copyright-and-Artificial-Intelligence-Part-2-Copyrightability-Report.pdf>

⁴⁵ <https://www.wired.com/story/ai-copyright-case-tracker/>

⁴⁶ https://en.wikipedia.org/wiki/Three_Laws_of_Robotics

⁴⁷ https://www.streetdirectory.com/travel_guide/120083/technology/the_fourth_law_of_robotics.html

	CN-AISGF2	US NIST-AIRMF 2022
Open-source model deployment	Requires registration, provenance tracing, and risk classification. High-risk models (e.g., reasoning agents) face stricter controls.	No registration required. Encourages internal risk mapping and documentation. Focus on transparency (see point 8) and reproducibility ⁴⁸
Enterprise use of generative AI	Content risk is paramount. Enterprises must implement red-teaming adversarial testing to uncover “ <i>unsafe behaviors</i> ”, watermarking in AI-generated content, and output “ <i>filtering</i> ”.	Emphasizes context-specific risk mitigation. Suggests bias audits (see point 8 and point 11 for our concerns), explainability, ⁴⁹ and human-in-the-loop safeguards (increasing algos interact with algos activities make this hard).
Cross-border AI collaboration	Subject to national security review. Export controls may apply to foundational models.	No formal restrictions, except under export control regimes.
	OECD ⁵⁰ and ISO (22989 terminology, 42001 management system, 42005 assessment, 42006 audit certification, 23894 Risk, 38507 governance, 5338 lifecycle, 23053 using ML, 24028 trustworthiness, 24027 bias, 5469 functional safety) ⁵¹ standards. Amid US past administration encouraged alignment with these standards, which we believe CN-AISGF2 has adopted. We have reservations with these long lists of “ <i>global AI requirements</i> ” may NOT fit the US best interest or are barriers affecting US AI firms’ competitiveness in overseas. OECD’s definition – “ <i>An AI system is a machine-based system that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments</i> ” seems overly focus on LLMs and undermines other cognitive reasoning AIs. One-size-fits-all global standards or lack sector tailoring is against big trends towards mass customization that unleash tremendous values. Prescribing the wrong framework undermines true AI risks and may inadvertently exacerbate the risks to humankind.	
Incident response (e.g., model hallucination)	May trigger risk reclassification. Requires reporting to National Computer Network Emergency Response Technical Team/Coordination Center of China. Hallucinations or “ <i>unsafe outputs</i> ” = • Loss of control over reasoning agents • Misuse / repurpose of open-source models for prohibited applications • Adversarial testing reveals “<i>vulnerabilities</i>” that were not mitigated	Suggests post-incident review, documentation, and updates to risk posture. No mandatory reporting unless regulated by sectoral law. Nemil Dalal argued that “ <i>today’s biggest threat to democracy isn’t fake news — it’s selective facts.</i> ” ⁵² AI hallucinations may discover unknown unknowns which were previously nonsensical to human. It is a paradigm shift to go from suspicions to opportunistic about newfound signals among chaos.

⁴⁸ “*Reproducibility of outcomes*” is related to “*consensus building*”, while trying to “*neutralize*” biases can dilute strategic advantage, especially when foreign adversaries are not playing by the same rules (see [point 8](#)).

⁴⁹ The foundation of a responsible AI is NOT about how good one person can articulate or reveal the secret ingredients of an AI to others. Unfortunately, many existing financial regulations are based on big law or consulting firms help their elite clients to brag about their policies and procedures to pursue enforcement, rather than reviewing real substance in improving controls. The related compliance burdens are ways for Elites to suppress smaller competitors who cannot afford to pay the big law or consulting firms. In case something bad happened, Elites leverage on the clout of big law or consulting firms to shoulder the blame. Small AI firms should allow to compete with larger counterparts. We despise policies exacerbating gap between the “*Haves*” and “*Have-Nots*”.

⁵⁰ <https://oecd.ai/en/ai-principles> ; <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>

⁵¹ <https://www.iso.org/sectors/it-technologies/ai> ;

⁵² <https://qz.com/1130094/todays-biggest-threat-to-democracy-isnt-fake-news-its-selective-facts>

11. Bias can be “conscious” and “unconscious”. A cognitive system does NOT have to be conscious. Neuroscientists believe consciousness could be a distributed process that does not depend on a singular “self”. Unconscious bias can inflict harm amid unintentional – negligence. Cognitive systems with no recommendations⁵³ should NOT escape AI responsibilities. “Externalities” arise from poorly defined private rights and high social costs, rather than being an inherent market failure.⁴¹ Enable clearer contracts between parties is better than rule-by-enforcement, unless there is a compelling case of market failure to intervene in rebuilding public trust in government.

Bias depends on social norm. Social norm evolves overtime. Hallucination is considered an output that is out of norms. Hallucinations are like dreams (a state of consciousness that one’s “awareness” of external environments may be out of synch), except dreams may be more vivid and emotion than hallucinations. Should cognitive systems be allowed to dream – a possible indicator of Artificial General Intelligence (AGI)? Again, AI hallucinations may discover unknown unknowns which were previously nonsensical to human.

Policy makers should encourage the industry to turn “unknowns” into “knowns.” It will enhance AI cognitive system’s performance to better understand nuances. We despise heavy-handed government policy to brutally force AI firms to censor / filtering so-called “unsafe behaviors / outputs” or require adversarial training to ban or reveal what authoritarian may constitute as “vulnerabilities”. Per point 8, trying to “neutralize” biases in pursuit of consensus or fairness can dilute the US strategic advantage, especially when foreign adversaries are not playing by the same rules.

12. We foresee increasing crossover activities between AI, tokenization and digital assets. There are innovative ways to detect illicit activity involving digital assets⁵⁴ and scrutinize any Repurpose or Reuse or Recycle of data. For example, Public Key Infrastructure⁵⁵ and uses a mix of digital traits to stitch and track provenance, monitor via Log-Chain, to validate against impersonation, verify and detect possible circumvention, and authenticate in case of compromise due to identity theft, account takeover, or abuse of access.

Feel free to contact us with any questions and please keep us posted where our expertise might be helpful.

Sincerely,

Kelvin To

Founder and President

Data Boiler Technologies, LLC

This letter is also available at: https://www.DataBoiler.com/index_hm_files/DataBoiler%20WHOSTP%2020251027R.pdf

Cc: Mr. David O. Sacks, Special Advisor for Artificial Intelligence and Cryptocurrency, The White House
The Honorable Rep. Aaron Bean, Rep. Pete Sessions, Rep. Blake Moore, Co-Chairmans of DOGE Caucus
Mr. Faisal D'Souza, Technical Coordinator, NITRD NCO, National Science Foundation
The Honorable Paul S. Atkins, Chairman of the SEC
The Honorable Hester M. Peirce, Commissioner of the SEC
The Honorable Caroline A. Crenshaw, Commissioner of the SEC
The Honorable Mark T. Uyeda, Commissioner of the SEC
Mr. Jamie Selway, Director, Division of Trading and Markets, SEC
Ms. Shira Perlmutter, Register of Copyrights and Director U.S. Copyright Office
Mr. Michael E. Clements, Director - Financial Markets and Community Investment, GAO

⁵³ <https://www.baeldung.com/cs/cognitive-computing-vs-ai>

⁵⁴ https://www.databoiler.com/index_hm_files/DataBoiler%20USDT%2020251017.pdf

⁵⁵ <https://www.linkedin.com/pulse/improving-trust-amid-race-technologies-kelvin-to-8vxrc>