

October 27, 2025

Co-authored by Melissa Hopkins, Ishan Mehta, Anita Cicero

RESPONSE TO AI REGULATORY REFORM RFI

Submitted by the Johns Hopkins Center for Health Security

Executive Summary

Thank you for the opportunity to provide comments in response to the Office of Science and Technology Policy's (OSTP) Request for Information (RFI) on Regulatory Reform on Artificial Intelligence.¹ The comments expressed herein reflect the thoughts of the Johns Hopkins Center for Health Security and do not necessarily reflect the views of Johns Hopkins University.

In order to expedite the development and deployment of AI-powered biotechnology and to alleviate barriers to AI development and deployment that arise from a lack of clarity or interpretive guidance on how existing rules cover AI activities, OSTP should:

- 1) Release the updated Framework for Nucleic Acid Synthesis Screening; and**
- 2) Clarify liability for biological harms by establishing an accreditation and certification program at the National Institute of Standards and Technology (NIST).**

Introduction

The Johns Hopkins Center for Health Security (CHS) conducts research on new policy approaches, scientific advances, and technological innovations that can strengthen health security and save lives. CHS has 25 years of experience in biosecurity and is dedicated to ensuring a future in which biological weapons in the hands of our adversaries can no longer threaten our world. CHS is composed of researchers and experts in science, national security, emerging technology, economics, law, medicine, and public health.

We are excited and optimistic about US leadership in leveraging AI to prevent and cure diseases, discover new life-saving medical products, improve public health, and generally improve the lives and livelihoods of citizens. AI technology also has tremendous potential to enhance both our economic well-being and our nation's geopolitical position. The next few years are critical, and we agree that it is advisable to avoid excessive regulations that attempt to eliminate all potential risks. Rather, it makes more sense to promote AI development and deployment in the public and private sectors while preventing foreign adversaries or other malicious actors from misusing our AI systems to create high-consequence chemical, biological, radiological, and nuclear (CBRN) weapons that could threaten America's national security interests. For biological threats in particular, the current threat landscape includes both deliberate use and high-consequence accidents by state actors and non-state actors such as terrorist groups or lone wolves.

¹ Office of Science and Technology Policy, *Notice of Request for Information; Regulatory Reform on Artificial Intelligence*, 90 Fed. Reg. 46422 (Sept. 26, 2025), <https://www.federalregister.gov/documents/2025/09/26/2025-18737/notice-of-request-for-information-regulatory-reform-on-artificial-intelligence>.

The focus of our comments here is specifically on regulatory reform to catalyze the development of AI systems trained in whole or in part on biological data while also scrupulously preventing the misuse of AI systems to develop high- consequence biological weapons.

Response

- (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?**

Release the Updated Framework for Nucleic Acid Synthesis Screening

The application of AI for biotechnology holds tremendous promise for American prosperity and security through applications like drug discovery, biomanufacturing, and medical imaging.²

One key component of realizing these benefits is the generation of synthetic nucleic acids for life science labs and researchers. However, a small subset of this technology is inherently dual-use—some sequences of synthetic nucleic acids can be used to create both harmful and beneficial biological substances, with future AI models or systems potentially capable of both significantly enabling non-experts to create and carry out biological weapons attacks as well as raising the ceiling of harm associated with such attacks.

We cannot allow AI to be misused as a bioweapon—it would not only severely damage American economic, health, and national security, but could also potentially slow or pause altogether the development of the beneficial uses of AI for biotechnology.

The American AI and biotech industries already lead the world in biosecurity, and the Administration can help to ensure that these standards are exported. Many frontier AI companies have publicly recognized that AI systems could increase the risk of biological attacks if used inappropriately by malicious actors with the ability to acquire dangerous materials and have implemented voluntary mitigation measures and invested considerable resources to address these risks within their own models.³ Many major nucleic acid synthesis providers actively support enhanced oversight policies, with more than 30 leading providers (mostly US companies) committing to voluntary nucleic acid synthesis screening.⁴ While these industry norms are a positive step, voluntary standards are not a replacement for federal guidance, which can level the playing field as compared to the status quo—which penalizes responsible actors who invest in safety, while others go scot-free.

² See NATIONAL SECURITY COMMISSION ON BIOTECHNOLOGY, CHARTING THE FUTURE OF BIOTECHNOLOGY (April 2025), <https://www.biotech.senate.gov/final-report/chapters/>.

³ See, e.g., Frontier Model Forum, *Issue Brief: Frontier AI Biosafety Thresholds* (May 12, 2025), <https://www.frontiermodelforum.org/issue-briefs/frontier-ai-biosafety-thresholds/>.

⁴ See International Gene Synthesis Consortium, <https://genesynthesisconsortium.org/>; see also Gene Synthesis Information Hub, *List of Framework-Attesting Nucleic Acid Synthesis Providers & Benchtop Manufacturers*, <https://genesynthesiscreening.centerforhealthsecurity.org/for-customers/list-of-framework-attesting-providers-benchtop-manufacturers>.

The US Framework for Nucleic Acid Synthesis Screening, last revised in September of 2024, established that US federal funding agencies will require that procurement of synthetic nucleic acids and benchtop nucleic acid synthesis equipment using federal life sciences funding be conducted through providers and benchtop equipment manufacturers that adhere to the Framework.⁵ The Framework incorporates and supplements parts of the Screening Framework Guidance for Providers and Users of Synthetic Nucleic Acids⁶ and its companion guide.⁷ Parts of the Framework affecting providers should have gone into effect April 26, 2025; however, the Framework was rescinded in January 2025.

The Executive Order on Improving the Safety and Security of Biological Research (May 5, 2025)⁸ calls for a revision or replacement of the Framework within 90 days. The EO still calls for agencies to tie Framework compliance to federal funding. It also calls for actions to expand the Framework to non-federally funded entities and requires the inclusion of enforcement mechanisms for non-compliance with the Framework. To date, no revision or replacement of the Framework has taken place despite the passing of the 90-day deadline.

While the Framework is delayed, providers are left with uncertainty without pertinent guidance from the federal government. The 2023 HHS Guidance is assumed to be in effect, but the May 2025 EO did not clarify if that was the case. Moreover, up-to-date guidance is needed that reflects the progress in current technology. By establishing clear federal requirements tied to US funding, the Framework would create a competitive advantage for compliant providers and pressure international competitors to meet these higher biosecurity standards to access the lucrative US market. The Framework would demonstrate US leadership in biosecurity and help export American safety standards globally, as many US companies are already voluntarily participating and setting the global standard for responsible nucleic acid synthesis screening.

By releasing the updated Framework, the White House can provide clarity to U.S. companies and providers as well as provide assurance to the leading AI labs that downstream biological risks are mitigated.

Clarify Liability for Biological Harms by Establishing an Accreditation and Certification Program at NIST

Section 230 of the Communications Decency Act (47 USC 230) states, “No provider . . . of an interactive computer service shall be treated as the publisher or speaker of any information *provided by another information content provider*” [emphasis added]. Section 230 effectively shields hosts of user-generated content from liability resulting from harms that the content generated by third parties on their platforms may cause. This shield has been central to preventing or defending against claims against social media companies, with recent case law centering around the question of the extent to which content on a

⁵ THE WHITE HOUSE, FRAMEWORK FOR NUCLEIC ACID SYNTHESIS SCREENING (April 2024), <https://aspr.hhs.gov/S3/Documents/OSTP-Nucleic-Acid-Synthesis-Screening-Framework-Sep2024.pdf>.

⁶ DEPT. OF HEALTH & HUMAN SERVS., SCREENING FRAMEWORK GUIDANCE FOR PROVIDERS AND USERS OF SYNTHETIC NUCLEIC ACIDS (Oct. 2023), <https://aspr.hhs.gov/S3/Documents/SynNA-Guidance-2023.pdf>.

⁷ DEPT. OF HEALTH & HUMAN SERVS., COMPANION GUIDE TO ASSIST IN IMPLEMENTING THE RECOMMENDATIONS OF THE SCREENING FRAMEWORK GUIDANCE FOR PROVIDERS AND USERS OF SYNTHETIC NUCLEIC ACIDS (Oct. 2023), <https://aspr.hhs.gov/S3/Documents/SynNA-Companion-Guide-508.pdf>.

⁸ EXEC. OFFICE OF THE PRESIDENT, 90 Fed. Reg. 19611.

platform is *platform-generated* or *third-party generated*.⁹

However, dicta from Justice Neil Gorsuch in the 2023 case of *Gonzalez vs. Google* indicated that generative AI content would not qualify for Section 230 liability protection.¹⁰ Justice Gorsuch, in outlining the petitioner’s argument, distinguished between AI-generated recommendations such as YouTube recommendations or search results and other forms of generative AI such as chatbots, saying:

“As I take your argument, you think that the Ninth Circuit's Neutral Tools Rule is wrong because, in a post-algorithm world, artificial intelligence can generate some forms of content, even according to Neutral Rules. I mean, artificial intelligence generates poetry, it generates polemics today. That -- that would be content that goes beyond picking, choosing, analyzing, or digesting content. And that is not protected. Let's -- let's assume that's right, okay? Then I guess the question becomes, what do we do about YouTube's recommendations? And -- and as I see it, we have a few options. We could say that YouTube does generate its own content when it makes a recommendation, says up next. We could say no, that's more like picking and choosing. Or we could say the Ninth Circuit's Neutral Tools test was mistaken because, in some circumstances, even neutral tools, like algorithms, can generate through artificial intelligence forms of content and that the Ninth Circuit wasn't sensitive to that possibility and remand the case for it to consider that question. What's wrong with that?”¹¹

The Court ultimately declined to rule on the Section 230 issue and instead remanded the case to the Ninth Circuit to reconsider its decision in light of its ruling in *Twitter, Inc. v. Taamneh*,¹² which was decided on other grounds.

Clarifying liability for generative AI companies that includes a safe harbor for actions taken to mitigate potential high-consequence biological harms would enable AI developers and deployers to innovate and deploy without fear of putting their companies at risk. Without such clarity, companies would likely be subject to significant liability risk under state and federal tort law. For example, a RAND report found that AI developers may incur such liability under negligence, products liability, and public nuisance doctrines – especially if those companies fail to adhere to industry best practices regarding safety and security.¹³ Trade associations such as the Frontier Model Forum are developing a list of shared best

⁹ See, e.g., *Zeran v. America Online, Inc.*, 129 F.3d 327, 330–31 (4th Cir. 1997) (holding that Section 230 protects platforms from liability even when they have notice of problematic content and fail to remove it).

¹⁰ See 598 U.S. 617.

¹¹ Tech Policy Press, *Transcript: Gonzalez v. Google Oral Argument* (February 21, 2023), <https://www.techpolicy.press/transcript-gonzalez-v-google-oral-argument/>.

¹² 598 U.S. 471 (2023).

¹³ See Ramakrishnan, Ketan, Gregory Smith, and Conor Downey, *U.S. Tort Liability for Large-Scale Artificial Intelligence Damages: A Primer for Developers and Policymakers*, RAND CORP. (August 21, 2024), https://www.rand.org/pubs/research_reports/RRA3084-1.html.

practices,¹⁴ but it's unclear to what extent those best practices would (or should) apply to smaller companies, academics, or start-ups without the resources to implement them.

The System Cards from some of the Frontier Model Forum's member companies indicate a shared set of biological capability benchmarks that are costly to run,¹⁵ which may be cost prohibitive for smaller companies, academics, and start-ups to conduct. This inability to compete with industry best practices thus puts smaller American companies at liability risk if their model is used to generate step-by-step directions, convey tacit scientific knowledge that is otherwise hard to learn, or generate novel information (as a content creator) that is used to build a bioweapon.

On the state level, there have been quite a few attempts at clarifying liability for developers and deployers, but often in a way that does little to protect the public or the nation from AI-related, large-scale biological harms. For example, SB 813 in the California legislature provides AI developers with a rebuttable presumption of reasonable care against any personal injury or property damage claims resulting from their models' output if the AI developer was certified by a "multi-stakeholder research organization" at the time of the plaintiff's injuries.¹⁶ The concept of certification from a "multi-stakeholder research organization" serving as a liability shield for AI developers has been articulated and explored several times, most recently by a former senior policy advisor on AI and emerging technology for the Trump Administration,¹⁷ and before that by Anthropic's Jack Clark and a former senior policy advisor to OpenAI.¹⁸

We support the utilization of biosecurity evaluations by third-party evaluators.¹⁹ However, SB 813 and the concept of a "multi-stakeholder research organization" as outlined by its proponents incentivize forum shopping for the third party with the least-stringent requirements for certification and are overly broad as written.

As an alternative to the SB 813 approach and a patchwork of state legislation, and to help shield AI developers and deployers from tort liability, the Administration should work towards establishing an accreditation program for third-party biosecurity evaluators and work with Congress to preempt a narrow set of state AI laws, such as those dealing with matters of national security such as biosecurity. To accomplish this, the Administration

¹⁴ See, e.g., Frontier Model Forum, *Issue Brief: Frontier AI Biosafety Thresholds* (May 12, 2025), <https://www.frontiermodelforum.org/issue-briefs/frontier-ai-biosafety-thresholds/>.

¹⁵ See, e.g., *System Card: Claude Sonnet 4.5*, ANTHROPIC (Sept. 2025), <https://assets.anthropic.com/m/12f214efcc2f457a/original/Claude-Sonnet-4-5-System-Card.pdf>; *GPT-5 System Card*, OPENAI (Aug. 13, 2025), <https://cdn.openai.com/gpt-5-system-card.pdf>; *Gemini 2.5 Deep Think Model Card*, GOOGLE DEEPMIND (Aug. 1, 2025), <https://storage.googleapis.com/deepmind-media/Model-Cards/Gemini-2-5-Deep-Think-Model-Card.pdf>; see also *Grok 4 Model Card*, X.AI (Aug. 20, 2025), <https://data.x.ai/2025-08-20-grok-4-model-card.pdf> (demonstrating that non-member frontier AI developers are also tending to use this industry best practice).

¹⁶ S.B. 813, 2025-2026 Reg. Sess. (Cal. 2025) (introduced by Sen. McNerney, Feb. 21, 2025).

¹⁷ Dean W. Ball, *Putting private AI governance into action* *Putting Private AI Governance into Action*, HYPERDIMENSIONAL (Mar. 20, 2025), <https://www.hyperdimensional.co/p/putting-private-governance-into-action>.

¹⁸ Gillian K. Hadfield & Jack Clark, *Regulatory Markets: The Future of AI Governance*, ARXIV (April 25, 2023), <https://arxiv.org/abs/2304.04914>.

¹⁹ See Hopkins, Melissa, Anita Cicero, & Tom Inglesby, *Response to AI Action Plan Request for Comment*, JOHNS HOPKINS CTR. HEALTH SEC. (March, 2025) https://centerforhealthsecurity.org/sites/default/files/2025-04/Johns-Hopkins-Center-for-Health-Security-AI-Action-Plan-RFI-3.20.25_0.pdf.

needs to first direct the Center for AI Standards and Innovation (CAISI) to issue biosecurity standards, including capabilities evaluations and mitigations, that are developed through a public comment and stakeholder process. Then, the Administration should direct NIST's National Voluntary Laboratory Accreditation Program to accredit third parties with the ability to provide certifications to companies that are able to comply with the CAISI standards. Under common law, it is likely that such certifications would serve as strong evidence as complying with industry best practices and effectively shield companies from potential tort claims.

CAISI currently runs its own capability evaluations in house and has not issued any standards for what they hope to see companies do if their models develop concerning biological capabilities. This work is unscalable if the United States hopes to conduct such testing while quickly bringing many models to market, as the government is not equipped to handle biosecurity capability evaluations for a large number of models. Being unable to have one's model evaluated by CAISI in a timely manner could either delay or otherwise prolong product launches or put companies in a tough position to be exposed to liability risk if they choose not to have their models evaluated by CAISI due to such delays. Outsourcing biosecurity capability evaluations and mitigation standards to third parties alleviates this bottleneck and provides a basic standard of biosecurity that the public expects.

We appreciate the opportunity to provide these comments to OSTP and look forward to working together to ensure that AI-powered biotechnology can flourish while maintaining robust safeguards against biological risks.