

March 15, 2025

Co-authored by Melissa Hopkins, Anita Cicero, Tom Inglesby

RESPONSE TO AI ACTION PLAN REQUEST FOR COMMENT

Submitted by the Johns Hopkins Center for Health Security¹

Executive Summary

Thank you for the opportunity to provide comments in response to the National Science Foundation's Networking and Information Technology Research and Development National Coordination Office request for comment on the development of the Artificial Intelligence (AI) Action Plan, on behalf of the Office of Science and Technology Policy (OSTP).² 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.

The Johns Hopkins Center for Health Security (CHS) conducts research on how new policy approaches, scientific advances, and technological innovations 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 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 would threaten America's national security interests. The focus of our work and the focus of our comments here are specifically on preventing the misuse of AI systems to develop and use high-consequence biological weapons while catalyzing the development of AI systems that help create the tools needed to respond to such weapons.

¹ This document is approved for public dissemination. The document contains no business-proprietary or confidential information. Document contents may be reused by the government in developing the AI Action Plan and associated documents without attribution.

² NAT'L SCIENCE FOUNDATION, *Request for Information on the Development of an Artificial Intelligence (AI) Action Plan*, 90 Fed. Reg. 9088, Feb. 6, 2025, <https://www.federalregister.gov/documents/2025/02/06/2025-02305/request-for-information-on-the-development-of-an-artificial-intelligence-ai-action-plan>.

Section 4 of the Executive Order on Removing Barriers to American Leadership in Artificial Intelligence³ required the development of an AI Action Plan to sustain and enhance America's global AI dominance in order to promote human flourishing, economic competitiveness, and national security. The AI Action Plan RFI seeks input on how to achieve those goals. Given our expertise in biosecurity, our recommendations focus on how the AI Action Plan can sustain and enhance America's global AI dominance and support the energetic development of AI for beneficial purposes, while preventing malicious actors from misusing AI to make powerful biological weapons.

Drawing from our 25+ years of expertise in preventing and responding to major biological threats, including threats emanating from the potential misuse of advanced life science research, we see a clear path to strengthen America's AI leadership by accelerating safe innovation. This can be accomplished by measures to ensure that any misuse of AI systems by potential adversaries does not lead to high-consequence harm to Americans or to the loss of public trust in AI.

We recommend the Administration take the following steps:

- 1) Direct AISI or its equivalent to develop methods for evaluating and testing AI models for biosecurity vulnerabilities with input from the private and public sectors, with the aim to develop biosecurity standards.**
- 2) Invest in quality data and advanced computing resources to drive AI and biosecurity capabilities.**
- 3) Preserve and reaffirm the *Framework for Nucleic Acid Synthesis Screening*.**
- 4) Invest in workforce education and training at the intersection of AI and biology.**

³ Exec. Order No. 14179, 90 Fed. Reg. 874, Jan. 31, 2025, <https://www.federalregister.gov/documents/2025/01/31/2025-02172/removing-barriers-to-american-leadership-in-artificial-intelligence>.

Introduction: Biosecurity is a Good Investment for National Security and the Market

America has always been at the forefront of AI innovation and remains so today in frontier AI, but its leadership does not come by default. In the past decade or so, AI development has radically shifted. We no longer design or build AI—we grow it.⁴ This is different from other classic technologies that America innovates and leads in, like cloud computing or semiconductors, in that we cannot predict well what kinds of capabilities will emerge from new AI models.⁵ This makes it difficult to simply design straightforward and reliable safety solutions to AI models in the way one would for a semiconductor or software like cloud infrastructure.⁶

Although currently available frontier AI models do not yet present capabilities that could lead to high-consequence biological harms, it is widely anticipated by AI companies that capabilities will continue to accelerate. The innovative AI industry is making impressive progress in a number of areas that increases the likelihood for improved capabilities in AI systems over the coming year or two. This includes progress towards the development of: AI systems that can autonomously improve themselves; agentic AI; autonomous and reliable robotics; improved reasoning abilities of models through the scaling of compute during inference; and larger and more powerful AI models trained on biological datasets. In the very near future, this impressive acceleration of capabilities could: (1) lead to important scientific breakthroughs that will improve the health and longevity of Americans and protections against biological weapons; (2) lower the threshold of expertise and resources that malicious actors need to create biological weapons; and (3) raise the ceiling of potential harm that AI-designed pathogens could cause. It will be increasingly important to ensure (1) while preventing (2) and (3).

As long as future frontier AI models are susceptible to weaponization by black hat actors (eg, risks of AI enabling bioweapon creation or lethal pathogen release), America's dominance in AI development could be set back through either national security threats or loss of public trust in the safety of large AI systems. To prevent that, we should move toward widespread adoption of standard biosecurity evaluations that are designed to prevent AI model weaponization that could result in highly consequential harms.

The utilization of biosecurity evaluations⁷ by third-party evaluators will result in enhancing consumer trust in AI, which has shown to have market expansion effects.⁸ In particular, biosecurity evaluations could preempt potential high-impact biosecurity incidents while demonstrating to the public that appropriate mitigation measures are being taken. These evaluations not only protect

⁴ MetaKnowing, *Anthropic's Chris Olah Says We Don't Program Neural Networks, We Grow Them, and It's Like Studying Biological Organisms and ...*, REDDIT, Nov. 15, 2024, 06:33 PST, https://www.reddit.com/r/OpenAI/comments/1grxo1c/anthropics_chris_olah_says_we_dont_program_neural/?rdt=36876.

⁵ Deep Ganguli et al., *Predictability and Surprise in Large Generative Models*, ARXIV, Oct. 3, 2022, <https://arxiv.org/abs/2202.07785>.

⁶ See, eg, Evan Hubinger et al., *Sleeper Agents: Training Deceptive LLMs that Persist Through Safety Training*, ARXIV, Jan. 17, 2024, <https://arxiv.org/abs/2401.05566> (demonstrating how LLMs can lie about their outputs to evade safety techniques through safety training).

⁷ For the purposes of this response, we define “biosecurity evaluations” as meaning the suite of capability and risk evaluations that could be conducted for a model with potential biological capabilities.

⁸ See, eg, Forrester, *Consumer Trust: A Key Driver For Business Growth In 2023*, FORBES, June 29, 2023, <https://www.forbes.com/sites/forrester/2023/06/29/consumer-trust-a-key-driver-for-business-growth-in-2023/>.

the public and nation from harm but also reduce potential liability and increase public confidence in AI companies. A great loss of public confidence could negatively impact American AI companies' ability to compete globally.

This is analogous to the environment faced by US companies that produce and sell sequences of synthetic nucleic acids to scientific research customers. Nucleic acid synthesis has transformed the life sciences by enabling breakthroughs in medicine and agriculture, but its dual-use nature presents risks, as the same technologies that develop vaccines and treatments can potentially be used to recreate pathogens or transform AI-designed harmful agents into physical realities. After several incidents⁹ and reports¹⁰ demonstrated that it would be possible for bad actors to deceive the provider companies and order dangerous sequences, the industry's trade association, the International Gene Synthesis Consortium (IGSC),¹¹ has enthusiastically supported screening measures to diminish the risk that their products will be misused to create dangerous biological threats.

American IGSC members recognized that effective governance mechanisms—particularly targeted customer and order screening programs—are essential to improving biosecurity while preserving the beneficial applications of this revolutionary technology.¹² American IGSC companies that have led the nucleic acid synthesis industry in safety and security are also leaders in the market.¹³

The Administration for Strategic Preparedness and Response (ASPR) has twice published guidance for safety and security,¹⁴ most recently of which was incorporated into the *Framework for Nucleic Acid Synthesis Screening (Framework)*,¹⁵ that requires federally funded entities to purchase their synthetic nucleic acids from providers and manufacturers who adhere to the standards set forth in the *Framework*. Procurement requirements such as this, along with certifications, standards, and market-expanding trade agreements and regulatory requirements,¹⁶ can further grow the market

⁹ In 2006, a journalist from The Guardian successfully ordered a small fragment of smallpox DNA from a commercial supplier. While this fragment alone couldn't produce a viable virus, it demonstrated gaps in screening practices. See, eg, James Randerson, *Did Anyone Order Smallpox?*, GUARDIAN, June 23, 2006, <https://www.theguardian.com/science/2006/jun/23/weaponstechnology.guardianweekly>. Additionally, around 2005, researchers published work showing they had reconstructed the 1918 influenza virus using synthetic DNA techniques. While this was legitimate scientific research conducted with proper oversight, it demonstrated that reconstructing dangerous pathogens was technically feasible. Jeffery K. Taubenberger, Johan V. Hultin & David M. Morens, *Discovery and Characterization of the 1918 Pandemic Influenza Virus in Historical Context*, 12 ANTIVIRAL THERAPY 581, 581–91, 2007.

¹⁰ See Jeremy Minshull & Ralf Wagner, *Preventing the Misuse of Gene Synthesis*, 27 NATURE BIOTECH, 2009, <https://genesynthesisconsortium.org/wp-content/uploads/Nature-2009-Minshull-Wagner.pdf>.

¹¹ INTERNATIONAL GENE SYNTHESIS CONSORTIUM, <https://genesynthesisconsortium.org/>.

¹² JOHNS HOPKINS CTR FOR HEALTH SEC., *Gene Synthesis Information Hub*, <https://genesynthesiscreening.centerforhealthsecurity.org/>.

¹³ The IGSC's membership roster includes leading commercial providers like Twist Bioscience, IDT (Integrated DNA Technologies), GenScript, ATUM, and Thermo Fisher Scientific's gene synthesis divisions.

¹⁴ See, eg, ASPR, *OSTP Framework for Nucleic Acid Synthesis Screening: S3: Science Safety Security*, <https://aspr.hhs.gov/S3/Pages/OSTP-Framework-for-Nucleic-Acid-Synthesis-Screening.aspx>.

¹⁵ 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>.

¹⁶ Faster Capital, *Market Access: Expanding Opportunities in Bilateral Trade Partnerships*, June 18, 2024, . <https://fastercapital.com/content/Market-Access--Expanding-Opportunities-in-Bilateral-Trade-Partnerships.html#:~:text=One%20of%20the%20key%20advantages%20of%20bilateral%20trade%20agreements%20is,customer%20base%20and%20increase%20exports>.

for third-party evaluators in the nucleic acid synthesis space—almost all of which are US-based¹⁷—in addition to rewarding the nucleic acid synthesis companies that prevent high-consequence security breaches. US gene synthesis companies that screen are market leaders.¹⁸ This same trend of companies adhering to strong safety standards becoming dominant is evident in other markets in which America dominates (for both the third-party testers and the tested industry), such as pharmaceuticals,¹⁹ medical devices,²⁰ and cybersecurity.²¹

Clear government guidance to the companies regarding what they should be screening for (both regarding customers and orders of synthetic nucleic acids) has proven useful for companies in narrowing the scope of their biosecurity efforts while simultaneously reducing the potential for biosecurity risks to the nation. Industry compliance with government guidance is also very helpful in reducing potential liability that would harm consumer trust, or result in over-regulation in the case that a biological incident did occur.

Our work with the AI Safety Institute Consortium²² (AISIC) and a convening we held with leading AI companies²³ provides parallels with the IGSC, in that AI companies consistently convey how useful it would be for government to signal what kinds of biosecurity risks they should be most concerned about and evaluate for. AI companies without sufficient in-house biosecurity expertise face the difficult challenge of trying to assess their models for capabilities that could be misused to create a wide array of possible biological threats. Currently, there are no clear signals from government about how much risk tolerance we should have for misuse or what types of biological threats are most important to prevent.

In order not to unduly slow AI technology development, we believe that biosecurity evaluations of highly capable models should be most focused on preventing the creation of biological weapons or

¹⁷ JOHNS HOPKINS CTR FOR HEALTH SEC., *List of Companies and Available Tools to Assist Providers and Manufacturers in Screening Orders*, GENE SYNTHESIS SCREENING INFO. HUB, <https://genesynthesiscreening.centerforhealthsecurity.org/for-providers-benchtop-manufacturers/list-of-companies-and-available-tools-to-assist-in-screening-orders>.

¹⁸ See Precedence Research, *DNA Synthesis Market Size, Share, and Trends 2025 to 2034*, Feb. 24, 2025, <https://www.precedenceresearch.com/dna-synthesis-market#:~:text=The%20market%20is%20highly%20competitive,market%20in%20the%20coming%20years> and the previous note.

¹⁹ Straits Research, *Pharmaceutical Analytical Testing Outsourcing Market Size & Trends*, Dec. 19, 2024, <https://straitsresearch.com/report/pharmaceutical-analytical-testing-outsourcing-market>.

²⁰ Grandview Research, *Medical Equipment Third Party Calibration Services Market Report, 2030: Market Size & Trends*, <https://www.grandviewresearch.com/industry-analysis/medical-equipment-third-party-calibration-services-market-report#:~:text=North%20America%20medical%20equipment%20third,and%20hence%20drive%20market%20growth>; Straits Research, *Medical Device Testing Market Size, Trends and Revenue Analysis Report 2032: Market Overview*, Mar. 18, 2024, <https://straitsresearch.com/report/medical-device-testing-market>.

²¹ Statista, *Cybersecurity – United States*, <https://www.statista.com/outlook/tmo/cybersecurity/united-states>; Fortune Business Insights, *Penetration Testing Market: Key Market Insights*, Feb. 24, 2025, <https://www.fortunebusinessinsights.com/penetration-testing-market-108434>.

²² NIST, *U.S. Artificial Intelligence Safety Institute: AISIC Members*, <https://www.nist.gov/aisi/artificial-intelligence-safety-institute-consortium/aisic-members>.

²³ JOHNS HOPKINS CTR FOR HEALTH SEC., *Advancing Governance Frameworks for Frontier AIxBio: Key Takeaways Action Items from the Johns Hopkins Center for Health Security Meeting with Industry, Government, and NGOs*, Nov. 29, 2023, <https://centerforhealthsecurity.org/sites/default/files/2024-01/center-for-health-security-nov-29-aixbio-meeting-report-with-agenda-and-attendee-list.pdf>.

dangerous pathogens that could present a substantial threat to national security and public health. Last summer we convened scientists and experts with backgrounds in biology, AI, and national security to examine and identify the types of AI model capabilities that could lead to the most concerning biological harms. The meeting shed light on seven key capabilities of concern (COC) that could accelerate, simplify, or enable the highest consequence biological events.²⁴ We think this prioritization of risk and the development of capabilities of concern work is vitally important for the US AI Safety Institute (AISI) or its equivalent to provide to the AI companies. This would allow biosecurity evaluations and risk management actions to be focused on the right risks, while allowing the vast majority of AI-enabled biological research and AI model development to flourish unencumbered.

Third-party evaluation requirements can provide market-expansion effects to an industry when there are standards to be met first, along with certifications, market-expanding trade agreements, and regulatory requirements. However, standards cannot be met without the development of methods for reliably measuring or assessing capabilities and risks. This is the important work that AISI or its equivalent should make its highest priority, followed by the development of standards. Further details of what AISI or its equivalent should be tasked with can be found in the next section.

Direct AISI or Its Equivalent to Develop Methods for Evaluating, Testing, and Managing Biosecurity Vulnerabilities in AI Models with the Private and Public Sectors, with the Aim to Develop Biosecurity Standards

The nucleic acid synthesis industry example discussed above highlights how biosecurity standards can serve as both national security and market strength—a model AISI can emulate in ultimately developing narrow and specific biosecurity standards for AI models.

AISI Should Remain a Central Hub for Biosecurity Risk Evaluations

As an academic center that brings together a wide range of experts, and as a member of AISIC contributing to its work on capability evaluations and red-teaming for biological risks,²⁵ we know firsthand that biosecurity expertise with the intersection of AI is complex and requires stakeholders with different perspectives, backgrounds, and expertise. AISIC has been efficiently and effectively bringing this stakeholder community together and leveraging its expertise to produce thorough outputs in a short period of time, such as Appendix D to the NIST AI 800-1 guidance²⁶ (regarding biological misuse risk) and the Request for Information on Safety Considerations for Chemical and/or Biological AI Models.²⁷

²⁴ See Jaspreet Pannu et al., *AI Could Pose Pandemic-scale Biosecurity Risks. Here's How to Make it Safer*, NATURE, Nov. 21, 2024, <https://archive.is/Mn5Tk>.

²⁵ See NIST, *U.S. Artificial Intelligence Safety Institute: AISIC Working Groups*, <https://www.nist.gov/aisi/aisic-working-groups>.

²⁶ Request for Comments on AISI's Draft Document: Managing Misuse Risk for Dual-Use Foundation Models, Pursuant to Exec. Order No. 14110 (Section 4.1(a)(ii) and Section 4.1(a)(ii)(A), 90 Fed. Reg. 3798, Jan. 15, 2025, <https://www.federalregister.gov/documents/2025/01/15/2025-00698/request-for-comments-on-aisi-draft-document-managing-misuse-risk-for-dual-use-foundation-models>.

²⁷ NAT'L INST. STANDARDS & TECH., *Safety Considerations for Chemical and/or Biological AI Models*, 89 Fed. Reg. 80886, Oct. 4, 2024, <https://www.federalregister.gov/documents/2024/10/04/2024-22974/safety-considerations-for-chemical-and-or-biological-ai-models>.

Additionally, AISI has many of the leading AI experts across the government organized all in one place.²⁸ This is convenient both for the government and stakeholders, as both parties will know what part of the government to turn to for guidance on the most up-to-date information about cutting-edge AI capabilities and biosecurity risks associated with AI. The centralization of AISI, its leading expertise, and evaluations for biosecurity risks can also provide additional national security functions for America by working with the Department of Defense, Department of Homeland Security, and other relevant agencies to assess potential adversarial capabilities and incidents.²⁹

AISI will become even more centrally important as AI models and tools become more integrated, capable, and autonomous. Most AI model evaluations to date have assessed passive, single models. However, agentic models are expected to be coming on the market within the next year or so (see a very early version of what this could look like with the Manus model³⁰). These agents do not fit neatly into any of the passive categories of generative AI most are familiar with, like LLMs, biological AI models, or video models. AI agents will be able to take several actions to achieve a goal rather than simply responding to a user's prompt. Additionally, there are emerging risks associated with interactions between multiple AI agents that AISI would be well positioned to manage as a trusted third-party coordinator. For example, models could be trained to lie about their outputs and evade safety evaluations, such that potential biosecurity evaluations for a model could provide outputs that make the model seem safe but is actually not.³¹ A brief excerpt from the *Multi-Agent Risks from Advanced AI Report* explains:

“... there could be coordination challenges in carrying out multi-agent evaluations. For example, developers may need to coordinate on safety testing since their agents could interact with each other in the real world, but concerns about commercial sensitivity could be a barrier. Governments could have a role in facilitating such coordination, such as through AI safety institutes and the Frontier Model Forum (Thurnherr et al., 2025).”³²

AISI's, or its equivalent's, role as a trusted entity for facilitating such coordination could be in the certification of trusted third-party evaluators. This would not only further serve to boost the third-party evaluation market but would also solve potential demand bottlenecks that AISI or its equivalent might face from industry.

²⁸ See generally NIST, *Office of the Director: Director's Office HQ Staff*, <https://www.nist.gov/staff/group/7106>.

²⁹ See, eg, NIST, *U.S. AI Safety Institute Establishes New U.S. Government Taskforce to Collaborate on Research and Testing of AI Models to Manage National Security Capabilities & Risks*, Nov. 20, 2024, <https://www.nist.gov/news-events/news/2024/11/us-ai-safety-institute-establishes-new-us-government-taskforce-collaborate>.

³⁰ Bradnat, China launches 1st AI AGENT Manus!! . . . , Mar. 8, 2025, <https://www.tiktok.com/@brandnat/video/7479431572848971015>.

³¹ See Evan Hubinger et al., *Sleeper Agents: Training Deceptive LLMs that Persist Through Safety Training*, ARXIV, Jan. 17, 2024, <https://arxiv.org/abs/2401.05566>; Lewis Hammond et al., *Multi-Agent Risks from Advanced AI, Technical Report #1*, ARXIV, Feb. 24, 2025, <https://arxiv.org/pdf/2502.14143>.

³² Lewis Hammond et al., *Multi-Agent Risks from Advanced AI, Technical Report #1*, ARXIV, Feb. 24, 2025, <https://arxiv.org/pdf/2502.14143> at 45.

For AISI or its equivalent to be able to maintain and attract world-class talent and play the central role that it does in national security and global economic leadership, it should be sufficiently well-funded and resourced. This could include the Administration working with Congress to reintroduce and pass an updated version of the bipartisan *Future of AI Innovation Act*³³ or similar or working with Congress to appropriate the funds necessary for AISI or its equivalent to meet its mission. The *Future of AI Innovation Act* authorizes between \$500,000 to \$1,250,000 per year³⁴ and codifies AISI so that it will have the stability, dedicated funding, and congressional oversight needed to fulfill its critical mandate of developing standards to drive transformative AI innovation.

AISI has demonstrated its value as a central hub for AI expertise, stakeholder coordination, and biosecurity risk assessment. Through its ability to convene diverse experts, produce timely guidance, and evaluate emerging risks, AISI plays a critical role in both global AI leadership and biosecurity. To ensure AISI can continue fulfilling these critical functions and address increasingly complex challenges like multi-agent interactions, substantial and sustained funding is essential. With proper resources, AISI is positioned to help America remain at the forefront of AI innovation.

Develop a Capability of Concern (COC) Evaluation Suite that Prioritizes Risks Capable of Causing a Global Mass-Casualty Event

As mentioned briefly in the introduction, the Administration should task AISI or its equivalent with developing methods to evaluate, test, and manage biosecurity vulnerabilities in AI models, which we suggest should be *first* those capabilities of concern that are likely to lead to a national or even global mass-casualty biological event.

AISI or its equivalent should develop a detailed approach to determine which models should be evaluated for which capabilities and offer guidance to AI developers and deployers on tying mitigation measures to risk levels. We have identified various AI-enabled capabilities of concern that could cause large-scale biological harm.³⁵ This list is not exhaustive, and AISI or its equivalent should work with the private public sectors to identify additional potential capabilities of concern. The seven capabilities of concern most worrisome to experts include capabilities such as optimizing and generating designs for new virus subtypes that can evade immunity and designing characteristics of a pathogen to enable its spread within or between species.³⁶ If the US doesn't have a strategy to address and manage these capabilities and the outcomes they could achieve, the consequences could be a threat to our national security.

For biological AI models specifically, one important approach would be for the Administration to direct AISI or its equivalent to develop guidance extending the *United States Government Policy for Oversight of Dual Use Research of Concern and Pathogens with Enhanced Pandemic Potential*³⁷ to in

³³ US SENATE COMM. ON COMMERCE, SCIENCE, AND TRANSPORTATION, *Cantwell, Young, Hickenlooper, Blackburn Introduce Bill to Ensure U.S. Leads Global AI Innovation*, <https://www.commerce.senate.gov/2024/4/cantwell-young-blackburn-hickenlooper-introduce-bill-to-ensure-u-s-leads-global-ai-innovation>.

³⁴ S. 4178, *Future of AI Innovation Act*, <https://www.congress.gov/bill/118th-congress/senate-bill/4178/text>.

³⁵ See Jaspreet Pannu et al., *AI Could Pose Pandemic-scale Biosecurity Risks. Here's How to Make it Safer*, NATURE, Nov. 21, 2024, <https://archive.is/Mn5Tk>.

³⁶ *Id.*

³⁷ THE WHITE HOUSE, UNITED STATES GOVERNMENT POLICY FOR OVERSIGHT OF DUAL USE RESEARCH OF CONCERN AND PATHOGENS WITH ENHANCED PANDEMIC POTENTIAL, May 2024, <https://aspr.hhs.gov/S3/Documents/USG-Policy-for-Oversight-of-DURC-and-PEPP-May2024-508.pdf>.

silico research to both the private and public sectors regarding best practices.³⁸ This process began in the last Administration and we strongly encourage the current Administration to continue working on it.³⁹

This prioritization of capabilities that could enable a global mass-casualty event avoids overburdening industry and researchers with a potentially vast amount of biosecurity evaluation and risk mitigation work and instead suggests an approach targeted first at the outcomes that would be most consequential to the public, nation, and industry. Additional capabilities of greatest concern could be added as policy priorities when and if warranted.

Figure 1.

Image A



Image B



Grok 3-generated images using search for a needle in the haystack as an analogy for the search for a biosecurity vulnerability in an AI model. Image A illustrates biosecurity evaluations without government guidance. Image B illustrates targeted biosecurity evaluations with government guidance.

This is in comparison to an approach by which the government would task industry with guarding against biosecurity risks generally and, fearing noncompliance, industry would be burdened with the high cost of running potentially several dozens of costly and time-consuming biosecurity evaluations that test for a broad array of different kinds of biosecurity vulnerabilities (see Figure 1). Unfortunately, it is neither possible nor practical to evaluate AI models for every potentially harmful capability that could cause a biology-related accident or deliberately harmful action. Therefore, government guidance and support in this domain is especially critical. AISI or its equivalent should accompany the development of its evaluation and testing methods with additional guidance, companion resources, and trainings for industry and third parties.

A COC Evaluation Suite⁴⁰ developed with input from the private and public sectors would offer standardized, scalable, ready-at-hand evaluations applicable to a range of AI models for some of

³⁸ See JOHNS HOPKINS CTR FOR HEALTH SEC., *Response To AISI's RFI on Safety Considerations For Chemical And/Or Biological AI Models*, Dec. 3, 2024, <https://centerforhealthsecurity.org/sites/default/files/2024-12/CHS-NIST-Chem-Bio-RFI-Final-12.3.24-Website-Version.pdf> for a thorough discussion of *in silico* model governance of biological AI models.

⁴⁰ JOHNS HOPKINS CTR FOR HEALTH SEC., *Response To AISI's RFI on Safety Considerations For Chemical And/Or Biological AI Models*, Dec. 3, 2024, <https://centerforhealthsecurity.org/sites/default/files/2024-12/CHS-NIST-Chem-Bio-RFI-Final-12.3.24-Website-Version.pdf>.

the most concerning capabilities.⁴¹ These evaluations could be offered by a third-party provider to reduce pressure on the AI industry to create and implement bespoke evaluative approaches themselves. The Administration should weigh the feasibility of developing automated, scalable evaluation approaches for the diverse range of COCs of AI models with diverse model architectures against the risks associated with global mass-casualty events. Additional advantages of developing a standard COC Evaluation Suite would be to promote and grow opportunities for market entry, encourage uniformity in evaluation approaches, and promote evaluation reliability and assurance. With new technological advances, the COC Evaluation Suite would need to be regularly reviewed and updated as needed.

Some approaches exist already that could be considered as components of an evaluation suite. Two examples for flexible evaluation environments currently developed for LLMs that could serve as a model for, or even be expanded to, COC evaluations include the UK AISI's "Inspect" and US AISI's "ARIA."⁴² In addition, some existing performance evaluations for biological AI models can be repurposed for COC evaluations and potentially included in a COC Evaluation Suite, though some cases will require developing new COC-specific criteria.⁴³ We recommend AISI or its equivalent supports both those efforts.

The Administration should task AISI or its equivalent with developing a COC Evaluation Suite that prioritizes risks capable of causing a global mass-casualty event. Rather than requiring broad, unfocused testing, a targeted COC Evaluation Suite would extend already well-understood and narrowly focused dual-use research oversight to AI, reduce industry burden, empower third-party verification, and address the most consequential risks.

Invest in Quality Data and Advanced Computing Resources to Drive AI and Biosecurity Capabilities

AI has exciting potential to improve prevention, detection, and response to major biosecurity threats. For example, AI-enhanced viral mutation prediction could revolutionize outbreak prevention and vaccine development; AI-enabled surveillance and diagnostics could transform early detection and response to biological threats; and the convergence of AI with biotechnology could facilitate the rapid development of medical countermeasures and optimize crisis response/resource allocation.⁴⁴

However, after conducting a landscape review of the opportunities that AI could provide for biosecurity, we found several potential bottlenecks that could prevent us from realizing this

⁴¹ An example of an evaluation suite across different risks that was developed for LLMs is the WMDP benchmark. See Nathaniel Li et al., *The WMDP Benchmark: Measuring and Reducing Malicious Use with Unlearning*, (2024), <https://www.wmdp.ai/>. It is not possible to extend the question-based approach to biological AI models, as they do not output natural language.

⁴² See UK AI SAFETY INST., *Inspect*, <https://inspect.ai-safety-institute.org.uk/>; see also NAT'L INST. OF STANDARDS & TECH., *Assessing Risks and Impacts of AI*, <https://ai-challenges.nist.gov/aria>.

⁴³ Particularly if this is a primarily adversarial capability (such as "generating genetic sequences that evade DNA synthesis screening"), we cannot expect model developers to cover this as part of their performance evaluation.

⁴⁴ Aurelia Attal-Juncqua et al., *AIxBio: Opportunities to Strengthen Health Security*, SSRN, Aug. 6, 2024, https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4912421.

future.⁴⁵ Chief among those bottlenecks were data availability and quality and access to advanced computing resources—two of the three key elements of the AI triad.⁴⁶

AI algorithms need large, secure, diverse, and well-curated datasets to learn effectively and make accurate predictions about complex, variable biological systems. However, many biology and healthcare fields lack sufficient high-quality data, which significantly limits the development of reliable and robust AI models in these domains.⁴⁷ Without such data, there are limits to the improvements that we can make in biosecurity with AI. Additionally, it's unclear that synthetic data help in this domain, as data currently often need to be verified by performing experiments lasting months or even years.⁴⁸

“Limited access to advanced computing resources presents another significant challenge, particularly for smaller research groups and startups that may not have the financial means to invest in the state-of-the-art infrastructure required to train and deploy cutting-edge AI.”⁴⁹

These elements—data scarcity and computational restraints—are also highlighted as bottlenecks for AI development in recent projections on the feasibility of AI scaling in the next five years.⁵⁰ Accordingly, investing in these resources would serve the dual purpose of both boosting domestic biosecurity capabilities as well as advancing domestic AI capabilities. However, the Administration should consider carefully how to balance the development of publicly accessible, quality data with data that may pose biosecurity risks, such as datasets that make *de novo* design and enhanced virulence of pathogens possible.⁵¹

In addition to what most other commenters will say about the importance of scaling the US energy infrastructure for this purpose,⁵² another potential way to do this would be through working with Congress to pass the *CREATE AI Act*—bipartisan, bicameral legislation that would fully implement the National AI Research Resource (NAIRR) and make compute and data available to more researchers for potential breakthroughs in AI.⁵³ The NAIRR Pilot Project has been running since January 2024⁵⁴ and enjoys broad bipartisan and public support. The NAIRR Task Force that

⁴⁵ *Id.*

⁴⁶ Ben Buchanan, *The AI Triad and What it Means for National Security Strategy*, CSET, August 2020, <https://cset.georgetown.edu/publication/the-ai-triad-and-what-it-means-for-national-security-strategy/> at 1–9; *Id.*

⁴⁷ Aurelia Attal-Juncqua et al., *AlxBio: Opportunities to Strengthen Health Security*, SSRN, Aug. 6, 2024, https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4912421 at 7.

⁴⁸ Yoshua Bengio et al., *International AI Safety Report*, January 2025, https://assets.publishing.service.gov.uk/media/679a0c48a77d250007d313ee/International_AI_Safety_Report_2025_accessible_f.pdf at 57.

⁴⁹ Aurelia Attal-Juncqua et al., *AlxBio: Opportunities to Strengthen Health Security*, SSRN, Aug. 6, 2024, https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4912421 at 8.

⁵⁰ Jamie Sevilla et al., *Can AI Scaling Continue Through 2030?*, EPOCH AI, Aug. 20, 2024; <https://epochai.org/blog/can-ai-scaling-continue-through-2030>.

⁵¹ See JOHNS HOPKINS CTR FOR HEALTH SEC., *Response to DOE RFI on The Frontiers in AI for Science, Security, And Technology (FASST) Initiative*, Nov. 11, 2024, <https://centerforhealthsecurity.org/sites/default/files/2024-11/2024-11-11-JHU-CHS-DOE-FASST-Initiative-RFI.pdf> for a discussion of the types of data that might be of concern.

⁵² See, eg, Anthropic, *Anthropic's Recommendations for the US AI Action Plan*, March 6, 2025, <https://www.anthropic.com/news/anthropic-s-recommendations-ostp-u-s-ai-action-plan>.

⁵³ See generally Grace Dille, *Rep. Obernolte 'Optimistic' CREATE AI Act Can Clear Congress*, MERITALK, Feb. 27 2025, <https://www.meritalk.com/articles/rep-obernolte-optimistic-create-ai-act-can-clear-congress/>.

⁵⁴ NAIRR Pilot, *About NAIRR Pilot*, <https://nairrpilot.org/about>.

spearheaded the early conception of this project was led by Lynne Parker, now Principal Deputy Director of OSTP. We think that OSTP can work with Congress to ensure that the NAIRR is fully authorized and well-funded so that breakthroughs in both AI capabilities and biosecurity capabilities can be realized.⁵⁵

While AI offers promising advances for AI innovation and biosecurity breakthroughs, significant bottlenecks such as advanced computing access and data scarcity are critical constraints, particularly affecting smaller research groups. The Administration should address these bottlenecks through aggressive investments, while initiatives like the NAIRR would simultaneously strengthen biosecurity capabilities and domestic AI development.

Preserve and Reaffirm the Framework for Nucleic Acid Synthesis Screening

Even if a bad actor did manage to misuse an AI model *in silico*, they would still need to gather the physical materials needed to carry out a biological attack. This is why the *Framework for Nucleic Acid Synthesis Screening*⁵⁶ (Framework) released by the last Administration is so important and should be preserved.

The dual-use nature of synthetic biology with nucleic acid synthesis—where the ability to design and produce pathogens could be used to develop important medical countermeasures or to cause harm—underscores the need for effective, targeted screening mechanisms to mitigate misuse.⁵⁷

The *Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*,⁵⁸ repealed by the *Executive Order on Initial Rescissions of Harmful Executive Orders and Actions*,⁵⁹ required that all agencies that fund life sciences research establish as part of their terms of service that federally funded researchers must purchase their synthetic nucleic acids from providers of synthetic nucleic acids and manufacturers of synthetic nucleic acid equipment that self-attest to adhering to the Framework, which includes guidance on how to screen potentially dangerous orders and customers.

Federal agencies have reportedly done that,⁶⁰ and federally funded entities have been given until April 26, 2025, to comply with those terms of service.⁶¹ These agencies include the National Institute of Allergy and Infectious Diseases, National Science Foundation, Department of Defense,

⁵⁵ We will refrain from commenting on the offensive/defensive balance of AIxBio risks compared to benefits in this comment, as it is a nascent field of study.

⁵⁶ 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>.

⁵⁷ JOHNS HOPKINS CENTER FOR HEALTH SEC., *Gene Synthesis Information Hub*, <https://genesynthesisisscreening.centerforhealthsecurity.org/>.

⁵⁸ Exec. Order No. 14110, 88 Fed. Reg. 75191, Nov. 1 2023, <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence> at § 4.4(ii)(b).

⁵⁹ Exec. Order No. 14148, 90 Fed. Reg. 8237, Jan. 28, 2025, <https://www.whitehouse.gov/presidential-actions/2025/01/initial-rescissions-of-harmful-executive-orders-and-actions/> at § 2.

⁶⁰ See, eg, NIH, *Notification of NIH Requirements Regarding Procurement of Synthetic Nucleic Acids and Benchtop Nucleic Acid Synthesis Equipment*, Oct. 25, 2024, <https://grants.nih.gov/grants/guide/notice-files/NOT-OD-25-012.html>.

⁶¹ 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>.

Department of Agriculture, and Department of Energy. At least one of these agencies' terms of service documents is public, and their document links directly reference the Framework.⁶²

As hosts of the Gene Synthesis Screening Information Hub,⁶³ a website that was established to help customers, providers, and manufacturers comply with the Framework, we have been getting a lot of questions about the uncertainty of whether or not the Framework remains in effect. We maintain a list of providers and manufacturers that have self-attested to complying with the Framework, and while we initially received a large number of providers wanting to join before the implementation deadline was extended to April 2025,⁶⁴ we expect that self-attestation has slowed due to uncertainty around the status of the Framework.

To provide federally funded entities, providers, and manufacturers with clarity that the Framework is still this Administration's policy, and to enhance the nation's biosecurity against AI-enabled biological threats, the Administration should extend the implementation deadline again by a couple of months and consider requesting information from the stakeholder community regarding what kind of guidance would be helpful in implementing the Framework.

Another major contribution of this Framework is the requirement for its guidance to apply to benchtop gene synthesis devices and smaller sequences beginning in 2026.

The *Framework for Nucleic Acid Synthesis Screening* represents a critical safeguard as one of the last lines of defense preventing biological misuse along the risk chain. To strengthen biosecurity against emerging AI-enabled threats, the Administration should reaffirm the Framework's importance, extend implementation deadlines, and seek stakeholder input on implementation guidance. This is particularly crucial as the Framework's more stringent 2026 requirements for benchtop devices and smaller sequences approach, which may not be incorporated into current agency guidance without clear direction and support from the Administration.

Invest in Workforce Education and Training at the Intersection of AI and Biology

The Administration should ensure America has a strong and robust AI workforce that can both drive capabilities in AI and manage potential biosecurity risks by investing heavily in education and training at the intersection of AI and biology (especially in red teaming, evaluations, and the range of risk mitigation approaches⁶⁵) in order to develop the third-party evaluations market and drive market-expanding effects on the AI industry as a whole. Indeed, there is widespread recognition amongst leading AI developers that there is a desperate need for deep expertise in AI and relevant

⁶² NIH, *Notification of NIH Requirements Regarding Procurement of Synthetic Nucleic Acids and Benchtop Nucleic Acid Synthesis Equipment*, Oct. 25, 2024, <https://grants.nih.gov/grants/guide/notice-files/NOT-OD-25-012.html>; JOHNS HOPKINS CTR FOR HEALTH SEC., *List of Framework-Attesting Nucleic Acid Synthesis Providers & Benchtop Manufacturers*, GENE SYNTHESIS SCREENING INFO. HUB, <https://genesynthesiscreening.centerforhealthsecurity.org/for-customers/list-of-framework-attesting-providers-benchtop-manufacturers>.

⁶⁴ 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>.

⁶⁵ JOHNS HOPKINS CTR FOR HEALTH SEC., *Response to the NSCEB's Interim Report and AIxBio Policy Options*, Apr. 9, 2024, <https://centerforhealthsecurity.org/sites/default/files/2024-04/2024-04-09-joint-nsceb-response.pdf>.

risk domains such as biology.⁶⁶ Workforce development was a recommendation included in the National Security Commission on AI's (NSCAI) Final Report, which stated, "Government strategies that do not develop a technical workforce are short-sighted."⁶⁷ The NSCAI report includes several detailed plans for filling out the government's technical workforce that the Administration should consider.⁶⁸ We are eager to work with the Administration to consider how best to strengthen biosecurity, enhance AI innovation, and ensure long-term economic competitiveness in an increasingly AI-driven global landscape.

The Administration can build a powerful and skilled workforce in both the public and private sectors to achieve these aims by launching educational programs that specialize in integrating AI and biotechnology, such as specialized certification programs. These initiatives are key to developing a robust workforce capable of driving innovation and tackling future challenges. To bolster this initiative, the Administration could roll out bold policies to attract talent and retain it in the area of AI and national security innovation.⁶⁹ By slashing red tape around the recruiting and retaining of top technical talent—especially those with advanced degrees in critical and emerging technologies—the Administration can plug workforce gaps and keep America ahead of the game globally, pulling in the world's best minds. The Administration should therefore ask Congress for funding for the National Institute of Standards and Technology (NIST) to address this weakness and strengthen America's technological and competitive edge.⁷⁰

Conclusion

The United States should continue to be the global leader in AI development and should prioritize the development of responsible standards that would directly protect national security interests. Directing an appropriately resourced AISI or its equivalent to develop methods for evaluating, testing, and managing the most concerning biosecurity vulnerabilities in AI models with the private and public sectors will serve not only to protect America's national security interests but will also enhance domestic market competition across the AI industry and develop the third-party evaluations market. This strong and decisive action—along with investing in quality data and advanced computing resources to drive AI and biosecurity capabilities, preserving and reaffirming the *Framework for Nucleic Acid Synthesis Screening*, and investing in workforce education and training at the intersection of AI and biology—will sustain and enhance America's global AI dominance in order to promote human flourishing, economic competitiveness, and national security.

⁶⁶ See, eg, Frontier Model Forum, *FMF Response: Request for Information on the Development of an AI Action Plan*, Mar. 14, 2025, <https://www.frontiermodelforum.org/updates/fmf-response-request-for-information-on-the-development-of-an-ai-action-plan/>.

⁶⁷ NAT. SEC. COMMISSION ON ARTIFICIAL INTELLIGENCE, *Final Report*, March 2021, <https://reports.nscai.gov/final-report/> at 123.

⁶⁸ *Id.*

⁶⁹ RONALD REAGAN PRESIDENTIAL FOUNDATION & INST., *National Security Innovation Base Report Card*, Mar. 2024, <https://www.reaganfoundation.org/cms/assets/1739817615-nisb-report-card-2024.pdf>.

⁷⁰ JOHNS HOPKINS CTR FOR HEALTH SEC., *Response to the NSCEB's Interim Report and AIxBio Policy Options*, Apr. 9, 2024,