



March 15, 2025

Office of Science and Technology Policy
Executive Office of the President
1650 Pennsylvania Avenue NW
Washington, DC 20502

Re: Developing a Flexible, Innovation-Focused U.S. Approach to AI Governance¹

Ryan Nabil
Director and Senior Fellow, Technology Policy
National Taxpayers Union Foundation
122 C St NW
Washington, DC 20001

On behalf of National Taxpayers Union Foundation (NTUF), I welcome the opportunity to submit the following comments in response to the White House's request for comments on the U.S. approach to AI regulation.² Located in Washington, DC, the National Taxpayers Union is the oldest taxpayer advocacy organization in the United States. Its affiliated think-tank, NTUF, conducts research on economic and technology policy issues of interest to taxpayers, including U.S. and international approaches to artificial intelligence, emerging technologies, and data protection.

NTUF appreciates the Trump Administration's recognition of the need to create a more favorable regulatory environment where artificial intelligence and AI-enabled business models can thrive and promote economic growth and technological innovation. As the White House seeks to develop the U.S. approach to AI in greater detail, it can strengthen the U.S. position as a global center of AI innovation. To accomplish that goal, the U.S. government needs to adopt a flexible, evidence-based approach to AI governance, distinguishing between widely varying applications of AI in different contexts and designing proportionate and context-specific rules accordingly. We believe that the U.S. national AI strategy would benefit from the following recommendations:

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 National Science Foundation, "Request for Information: National Priorities for Artificial Intelligence," Federal Register 90, no. 24 (February 6, 2026): 9088, <https://www.federalregister.gov/documents/2025/02/06/2025-02305/request-for-information-on-the-development-of-an-artificial-intelligence-ai-action-plan>

1. The United States needs to adopt a flexible, innovation-focused approach that outlines the government's AI principles, establishes the U.S. AI framework, creates mechanisms to implement it, and develops measures to promote innovation and mitigate AI risks.
2. The United States would benefit from more closely evaluating the AI governance strategies of major jurisdictions—like the European Union, the United Kingdom, Japan, and Switzerland—in understanding how best to design a flexible, well-balanced approach to AI.
3. Given the widely divergent applications of AI to different sectors and business functions, the U.S. should regulate the applications of AI, rather than the underlying technology.
4. Well-designed AI sandbox programs can help improve the regulatory understanding of AI technologies and business models, design more flexible AI rules, and promote innovation.
5. Designing reciprocal sandbox arrangements with like-minded jurisdictions—such as the UK, the EU, and Switzerland— can promote cross-border innovation and regulatory cooperation.
6. The U.S. government should strengthen bilateral cooperation with like-minded partner countries and contribute more actively to developing international AI norms through multilateral institutions, such as the Organisation for Economic Cooperation and Development (OECD) and the Global Partnership for AI.

I. Developing a Flexible, Innovation-Focused Approach to AI Governance

While the U.S. federal government has rightly avoided passing a one-size-fits-all regulatory framework for AI, it has also lagged in developing a flexible, carefully calibrated, and evidence-based approach to AI governance. Yet, as state legislatures have sought to pass legislation related to AI applications, the United States faces the potential risk of growing regulatory fragmentation at the federal and state levels. Such a development is especially likely if the Trump administration and Republican lawmakers apply an overly binary approach to AI governance—with any AI-related legislation and regulations being considered “bad” and “harmful” and the absence of such regulation being perceived as a positive development without exception.

However, such binary thinking is unlikely to help the United States grapple with the complex legal and distinct challenges associated with various context- and sector-specific artificial intelligence applications. While an overly restrictive federal AI framework would threaten U.S. innovation—as noted in NTU’s filing to the Biden administration—the absence of a coherent approach could also heighten the risk of AI misuse and result in an increasingly complex patchwork of state regulations. Such a development would lead to a more fractured U.S. digital economy, hindering technological innovation and economic growth.³

Therefore, while the United States should refrain from passing one-size-fits-all comprehensive AI legislation that could constrain regulatory flexibility and struggle to keep pace with technological change and emerging risks, it should seek to create a flexible, principles-based AI framework that develops well-calibrated and proportionate rules according to the specific risks associated with AI use

³ Ryan Nabil, “Letter to the White House: The Need for A Flexible and Innovative AI Framework,” July 7, 2023, <https://www.ntu.org/foundation/detail/letter-to-the-white-house-the-need-for-a-flexible-and-innovative-ai-framework>.

in a specific context. Without a well-balanced, carefully designed regulatory strategy, the United States runs the risk of hampering the country's long-term AI potential.

In developing the national AI framework, U.S. lawmakers would benefit from evaluating the AI governance approaches of leading jurisdictions such as the EU, the UK, and Japan. While a detailed discussion of such national strategies goes beyond the scope of this submission, understanding different regulatory approaches—particularly between the EU and the UK—can be instructive in designing a flexible, evidence-based approach to AI governance.

As is the case under many civil law jurisdictions, the EU's approach to AI regulation is characterized by detailed and carefully negotiated legislation that seeks to predict and mitigate future risks from AI applications—as opposed to developing broader statutory principles and enabling regulators and courts to play a more active role in determining how such principles should apply to specific AI applications in light of new technological developments.

Last year, the European Parliament, the European Union's legislative organ, passed the Artificial Intelligence Act, the world's first comprehensive AI legislation to regulate AI use in almost every sector across the European single market.⁴ Under EU constitutional law, certain legislation like the AI Act require approval by a qualified majority (i.e., at least 15 out of 27 EU Member States) in the Council of the EU and then a simple majority in Parliament—a process often resulting in multiple rounds of negotiations and redrafting before the proposed legislation is ultimately approved. Therefore, the procedural benefits of passing single comprehensive legislation instead of multiple sectoral laws are all too understandable in the European context. Nevertheless, some of the AI Act's restrictive proposals, such as its vague and overly broad definition of AI and classifications of high-risk AI systems, risk hampering Europe's innovation potential, as pointed out by leading European scientists and policymakers,⁵ numerous companies like Siemens and private-sector bodies such as the German AI Association,⁶ as well as national and regional governments.⁷

⁴ European Commission, "Proposal for a Regulation of the European Parliament and of the Council Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts," COM (2021) 206 final (April 21, 2021), <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex%3A52021PC0206>.

⁵ Patrick Glauner and Kai Zenner, "KI-Verordnung – Bärendienst für die heimischen KMU" ["AI Regulation - Disservice to Domestic SMEs"], *Der Tagesspiegel*, April 19, 2023, <https://background.tagesspiegel.de/digitalisierung/ki-verordnung-baerendienst-fuer-die-heimischen-kmu>.

⁶ KI-Bundesverband [German AI Association], "Positionspapier des KI-Bundesverband e.V. zur EU-Regulierung von Künstlicher Intelligenz" ["Position Paper of the German AI Association on the EU's AI Act"], March 2021, https://ki-verband.de/wp-content/uploads/2022/02/KI_Regulierung_DE-komprimiert.pdf. Javier Espinoza, "European companies sound alarm over draft AI law," *Financial Times*, June 30, 2023, <https://www.ft.com/content/9b72a5f4-a6d8-41aa-95b8-c75f0bc92465>.

⁷ Benoit Berthelot, "Macron Calls for French AI Innovation as EU Votes to Regulate," *Bloomberg*, June 14, 2023, <https://www.bloomberg.com/news/articles/2023-06-14/macron-calls-for-french-ai-innovation-after-eu-votes-for-ai-act-restrictions>. Bayerische Staatsregierung [Bavarian State Government], "Studie zu KI-Regulierung: EU-Regeln stellen Unternehmen vor große Hürden / Digitalministerin Gerlach: Innovation nicht durch Überregulierung ausbremsen" ["Study on AI Regulation: EU Rules Will Pose Major Hurdles for Companies/Digital Minister Gerlach: Do not slow down innovation through overregulation"], press release, March 28, 2023, <https://www.bayern.de/studie-zu-ki-regulierung-eu-regeln-stellen-unternehmen-vor-grosse-huerden-digitalministerin-gerlach-innovation-nicht-durch-ueberregulierung-ausbremsen/>.

In contrast, the UK has advocated a more flexible, context-specific approach to AI, which seeks to regulate AI applications in different contexts rather than the underlying AI technologies. Instead of developing comprehensive AI legislation like the EU's AI Act, the UK government has proposed AI principles and a non-statutory AI framework, which regulators would apply to AI applications within their remit.⁸ Case law and jurisprudence by English courts would further clarify how existing statutes apply to AI applications, and the government reserves the right to introduce legislation to update sectoral rules if and when necessary.

Like the UK, the Japanese government has also advocated a light-touch, principles-based approach to AI regulation, which aims to promote innovation and economic growth in light of Japan's economic and demographic challenges.⁹

Given the similarity of the English and U.S. legal systems, we believe the UK's flexible, pro-innovation approach represents a better-suited model for the U.S. than the EU's current approach to AI governance. A well-calibrated, context-specific approach would allow the United States to remain flexible in updating its regulatory frameworks in light of new technological developments and emerging risks. Such an approach would also make it easier for sectoral legal frameworks to remain technology-neutral and allow regulators to apply the same rules and standards to the application of other emerging technologies like quantum computing and communications—instead of having to develop new legal frameworks and enact separate statutes for each new wave of technologies.¹⁰ To that end, the U.S. government should consider designing a flexible AI framework that outlines broader U.S. AI principles and guidelines for regulators and includes, amongst others, mechanisms to implement the AI framework and policies to encourage innovation and mitigate future risks.

II. Proportionate, Context-Specific Framework for Regulating AI in Different Sectors

The U.S. government should adopt a proportionate, context-specific approach to develop well-calibrated rules for different uses of AI technologies in various sectors. A major difference between AI and many previous technologies—such as atomic energy and space technologies—is AI's potential uses in a much wider segment of the economy, from healthcare to retail and financial services. The specific risks that AI poses in such sectors depend on the precise context in which AI is used rather than the underlying technologies themselves. Therefore, a proportionate approach to AI regulation should consider the precise context in which AI is used and develop well-calibrated rules

⁸ UK Department for Science, Innovation and Technology (DSIT) and the Office for Artificial Intelligence, "A Pro-Innovation Approach to AI Regulation," policy paper, updated June 22, 2023, <https://www.gov.uk/government/publications/ai-regulation-a-pro-innovation-approach/white-paper>.

⁹ Hiroki Habuka, "Japan's Approach to AI Regulation and Its Impact on the 2023 G7 Presidency," Center for Strategic and International Studies, February 14, 2023, <https://www.csis.org/analysis/japans-approach-ai-regulation-and-its-impact-2023-g7-presidency>. Ryan Morrison, "Japan becomes latest country proposing hands-off AI regulation, but businesses 'likely to follow EU rules,'" Tech Monitor, July 4, 2023, <https://techmonitor.ai/technology/ai-and-automation/japan-ai-europe-regulation-artificial-intelligence>.

¹⁰ Ryan Nabil, "Consultation Response to the UK Office for Artificial Intelligence: Principles for a Pro-Innovation Approach to AI Governance," National Taxpayers Union Foundation, June 21, 2023, <https://www.ntu.org/foundation/detail/consultation-response-to-the-uk-office-for-artificial-intelligence-principles-for-a-pro-innovation-approach-to-ai-governance/>

for specific uses instead of setting fixed rules and risk ratings for AI use in all sectors or even within the same sector.¹¹

For example, the use of AI chatbots for retail customer support are typically associated with less significant risks than AI applications in medical diagnostics and the healthcare sector. Accordingly, a context-specific, proportionate approach should consider the risks associated with AI applications in different circumstances and calibrate rules accordingly. Likewise, even within high-risk sectors, such as critical infrastructure, not all AI use poses the same level of risk. For instance, whereas using AI algorithms to optimize the operations of a nuclear plant carries significant risk, its use to detect minor cosmetic flaws, like surface damages, within the same plant typically carries much lower risks. Accordingly, classifying entire sectors as low or high-risk would not constitute a proportionate regulatory approach.¹²

Instead, a more sensible approach would entail the creation of a context-specific AI framework that sets out the overall AI principles and clarifies the regulatory characteristics of such a framework (Table A1). For example, the UK has adopted five AI principles based on the OECD's guidelines for trustworthy AI: i) Safety, security, and robustness; ii) Appropriate transparency and explainability; iii) Fairness; iv) Accountability and governance; and v) Contestability and redress.¹³ Likewise, the Japanese government—whose policy document contributed to formulating the OECD's AI principles—recognizes and suggests similarly phrased principles in its AI governance guidelines.¹⁴

Once the general principles are developed, they should form the basis of an overall AI framework. The framework should develop guidelines for sectoral regulators to apply the framework to specific AI uses in different contexts according to the specific risks they pose (Table A2). Regulators would then regulate AI within their remit while adhering to the guidelines outlined in the AI framework.

Ultimately, for such a framework to be practical in the U.S. context, Congress would need to provide a statutory basis for establishing U.S. AI principles, creating oversight over regulators for applying AI rules uniformly across different sectors, and developing mechanisms for inter-agency coordination. Furthermore, to ensure that sector-specific AI rules do not hamper innovation, U.S. lawmakers should also consider adding innovation as a statutory duty for regulators in enforcing the AI framework. Such a measure would help ensure that regulators not only consider identified and prioritized AI risks in agency rulemaking but that they also consider the potential risks of slowed innovation due to an overly restrictive regulatory approach.¹⁵

¹¹ Nabil, "UK Approach to AI Governance." DSIT, "Pro-Innovation Approach to AI."

¹² Ibid.

¹³ Ibid. Note that the OECD's principles are worded slightly differently: i) "inclusive growth, sustainable development, and well-being"; ii) "human-centred values and fairness"; iii) "transparency and explainability"; iv) "robustness, security and safety"; and v) "accountability". Organisation for Economic Co-operation and Development, "OECD AI Principles Overview," n.d., <https://oecd.ai/en/ai-principles>.

¹⁴ Japanese Ministry of Economy, Trade, and Industry (METI), Expert Group on How AI Principles Should Be Implemented, "Governance Guidelines for Implementation of AI Principle," January 28, 2022, https://www.meti.go.jp/shingikai/mono_info_service/ai_shakai_jisso/pdf/20220128_2.pdf.

¹⁵ Nabil, "UK Approach to AI Governance." DSIT, "Pro-Innovation Approach to AI."

III. Mechanisms to Support the Implementation of the U.S. AI Framework

The U.S. government should consider developing mechanisms to support the implementation of the U.S. AI framework and help ensure that AI principles and guidelines are applied uniformly across different sectors. While a principles-based, context-specific approach to AI would allow the United States to develop flexible and well-calibrated rules for AI in different sectors, this strategy comes with certain challenges that would need to be addressed in the U.S. AI framework.

A central challenge is that, since individual regulators have the flexibility to issue guidelines and adjust rules based on broader AI principles, there is a risk that such guidelines are not applied uniformly across different sectors.¹⁶ Such differences would not only create market uncertainties but would also pose a particular challenge when certain AI applications come under the jurisdiction of multiple regulators. A hypothetical example would entail the regulation of an AI-enabled investment advisory product dealing with the personal data of users—which could be subject to the overlapping jurisdiction of the Federal Trade Commission, the Federal Deposit Insurance Corporation, the Consumer Financial Protection Bureau, and even state regulators.¹⁷ The U.S. AI framework should, therefore, design preemptive mechanisms to address the potential challenge of regulatory inconsistencies that could arise in a more flexible, decentralized AI governance approach.

The UK’s proposed mechanisms for the implementation of its AI framework could provide a useful starting point for U.S. policymakers for thinking more analytically about such issues and designing policies accordingly. The UK’s AI White Paper proposes seven supporting mechanisms for the following objectives: i) monitoring the overall effectiveness of the AI framework; ii) supporting the coherent application of AI principles across the economy; iii) assessing and addressing cross-sectoral risks from AI applications; iv) providing support and guidance to businesses; v) improving business awareness and consumer awareness of trustworthy AI; vi) conducting horizontal scanning for emerging risks and regulatory trends; and vii) monitoring global regulatory developments.¹⁸

While the precise mechanisms would need to be calibrated and adapted to U.S. policy objectives and regulatory architecture, these proposals point to important challenges that U.S. lawmakers should consider while pursuing a more decentralized approach to AI regulation. The table below provides some potential mechanisms—based on the UK government’s AI White Paper—that Congress and the Trump administration could consider while designing the U.S. AI framework (Table 1).

¹⁶ Nabil, “UK Approach to AI Governance.” DSIT, “Pro-Innovation Approach to AI.”

¹⁷ Ryan Nabil, “How Regulatory Sandbox Programs Can Promote Technological Innovation and Consumer Welfare: Insights from Federal and State Experience,” Competitive Enterprise Institute OnPoint, no. 281 (2022), <https://cei.org/studies/how-regulatory-sandbox-programs-can-promote-technological-innovation-and-consumer-welfare/>.

¹⁸ DSIT, “Pro-Innovation Approach to AI.”

Table 1. Functions to Support the Implementation of a Potential U.S. AI Framework

Functions	Potential Activities
1) Monitoring, Assessment, and Feedback	i) Develop and maintain monitoring and evaluation mechanisms to assess the economic impacts of the U.S. AI framework across different sectors and for the entire economy. ii) Collect data and stakeholder input from regulators, the private sector, think tanks, and academic institutions to evaluate the U.S. AI framework's overall effectiveness. iii) Monitor the framework's effectiveness in maintaining a proportionate approach. iv) Assess the effectiveness of interagency regulatory coordination.
2) Coherent Implementation of AI Principles	i) Develop guidelines to support regulators in implementing the U.S. AI framework. ii) Identify potential inconsistencies in the way different regulators apply AI principles. iii) Create a platform for regulators to discuss and address regulatory inconsistencies. iv) Monitor the continued relevance of the AI principles established in the framework.
3) Cross-Sectoral Risk Assessment	i) Create a risk register of potential AI risks to evaluate different risks and support the development of the cross-sector risk assessment framework. ii) Monitor and review prioritized risks and identify emerging risks. iii) Provide a platform to clarify regulatory responsibilities, issue joint regulatory guidance, and share regulatory best practices.
4) Support for Innovators	i) Identify potential regulatory barriers to AI innovation in different sectors. ii) Assist regulators in creating and monitoring the effectiveness of AI sandboxes.
5) Education and Awareness	i) Provide informal guidance to businesses on navigating the AI regulatory landscape. ii) Advise start-ups and companies on applying to the appropriate sandbox. iii) Improve consumer awareness and public trust about how AI is regulated in the U.S. iv) Support the creation of innovation hubs, which are typically launched by regulators to provide start-ups and companies information about AI-related legal obligations, identify business opportunities, and invest in the U.S. AI ecosystem. Innovation hubs can also help start-ups identify and apply to the appropriate sectoral AI sandbox.
6) Horizontal Regulatory Scanning	i) Monitor emerging trends in U.S. and international AI governance, new technological developments, and emerging AI risks. ii) Work with actors from the private sector, universities, and think tanks to identify, prioritize, and mitigate emerging risks.
7) International Regulatory Frameworks	i) Monitor AI-related foreign legislation and global regulatory developments and evaluate potential implications for the U.S. regulatory approach and the broader AI ecosystem. ii) Provide recommendations on improving cross-border regulatory cooperation on AI. iii) Monitor alignment between the U.S. and international AI frameworks developed by multilateral organizations like the OECD and the Global Partnership on AI. iv) Evaluate U.S. compatibility with global AI standards and identify opportunities to harmonize standards and reduce barriers to trade and cross-border data flows. iv) Recommend policies based on the successes and failures of regulatory approaches in the EU, the UK, Japan, and other major jurisdictions.

Source: Author based on recommendations by DSIT and Office for AI (2023).¹⁹

¹⁹ DSIT, "Pro-Innovation Approach to AI."

IV. Risk Assessment Mechanisms to Identify and Mitigate Future AI Risks

A major challenge in AI governance is to develop a proportionate risk-management framework to identify, prioritize, and mitigate potential risks. The differences in how various jurisdictions seek to evaluate and mitigate such risks can provide insights into how U.S. lawmakers could develop an agile, multi-stakeholder framework to identify and mitigate future risks. At the risk of oversimplification, the EU's AI Act classifies AI systems into four categories of risks i) "Minimal-risk" AI systems, which require AI developers to comply with a code of conduct; ii) "Limited-risk" AI systems that require providers to comply with certain transparency requirements; iii) "High-risk" AI systems that must undergo a more rigorous conformity assessment; and iv) AI systems with "unacceptable risks," which are banned across the EU (Table A3). The EU also provides lists of AI usage that would be classified as "limited" and "high risk." (Table A3).²⁰

While the European Union's risk-based approach sounds reasonable on a *prima facie* basis, it has two major problems. First, this approach does not provide a flexible framework that adequately distinguishes between risks associated with different AI applications within the same sector. For instance, whereas the EU's AI Act treats all AI-enabled tasks related to the operation and management of critical infrastructure as "high risk,"²¹ the UK government is more careful in recognizing that, even within high-risk sectors like critical infrastructure, not all AI-enabled tools carry the same risks and should not be subject to uniform compliance and liability standards.²²

Under the EU's AI Act, many low-risk AI applications within sectors classified as "high-risk"—such as education, employment, and law—are therefore potentially subject to significantly more restrictive regulations than under the UK's AI framework (Table A3). For example, since the act considers the use of AI in education as high risk, AI-enabled language proficiency examinations by online platforms—which often provide a much cheaper and more accessible alternative to traditional language proficiency tests like the TOEFL and IELTS—would be subject to the same compliance standards as the use of AI in other high-risk areas like medical diagnostics and critical infrastructure.²³ Such a restrictive approach risks hampering innovation in online learning platforms, legal services, and other areas that the EU's AI Act classifies as "high risk."²⁴

Notwithstanding the European Union's well-informed, detailed approach to AI governance, the AI Act's risk assessment framework might struggle to be flexible in addressing future risks. Although generative AI chatbots and applications have become widespread in the last two years, the pace and scope of their rapid development would have been difficult to predict even five years ago. Likewise,

²⁰ Lilian Edwards, "The EU AI Act: a summary of its significance and scope," Ada Lovelace Institute, April 2022, <https://www.adalovelaceinstitute.org/wp-content/uploads/2022/04/Expert-explainer-The-EU-AI-Act-11-April-2022.pdf>.

²¹ Dechert LLP, "European Commission's Proposed Regulation on Artificial Intelligence: Conducting a Conformity Assessment for High-Risk AI - Say What?," November 16, 2021, <https://www.dechert.com/knowledge/onpoint/2021/11/european-commission-s-proposed-regulation-on-artificial-intellig.html>.

²² DSIT, "Pro-Innovation Approach to AI."

²³ Ibid.

²⁴ Ryan Nabil, "The EU's Recently Proposed Artificial Intelligence Act Goes Too Far," The National Interest, August 21, 2021, <https://nationalinterest.org/blog/buzz/eu-s-recently-proposed-artificial-intelligence-act/C2%A0goes-too-far-191733>.

despite the best efforts of lawmakers, regulators, and technologists alike, the business of making predictions about future AI risks remains a highly uncertain one. As such, it is difficult to accurately predict the AI landscape ten years from now and the unique set of risks and challenges such developments will pose. In the U.S. context, where technology-related legislative activities often lacks the same deliberative, long-termist approach more characteristic of the EU and UK’s decision-making processes, adopting a similar approach of classifying prespecified AI uses as “high risk” in statute—based on a static understanding of risks based on the technological landscape today—risks creating a regulatory framework that is less adept in identifying and mitigating future AI risks.

The UK government’s proposed strategy of continuously monitoring AI risks and enabling public-private collaboration to identify emerging risks represents a more flexible approach to risk management—one that also merits close examination in the U.S. context (Table A4). Instead of classifying a list of AI applications as high risk, the government has proposed a principles-based risk assessment framework, which sectoral regulators will use to evaluate risks within their regulatory remit. Furthermore, the UK has proposed the creation of “central risk functions” — separate from sectoral regulators—that would play a central role in monitoring the effectiveness of the AI framework, monitoring current and future AI risks, and providing advice to the government on which risks should be prioritized.²⁵ With closer regulatory cooperation between the government, regulators, and the private sector, this approach is more likely to enable more robust monitoring of potential AI risks and introduce or calibrate appropriate statutory instruments to address risks as they emerge.²⁶

A comparable U.S. mechanism—involving Congress and the federal government, sectoral regulators, the private sector, scientific experts, and independent risk evaluators—could be designed to identify and respond to future AI risks (Table A4). As part of this arrangement, Congress and the federal government would establish the overall U.S. AI framework and clarify risk management guidelines for sectoral regulators based on the AI framework. In turn, the sectoral regulators would enforce such guidelines within their regulatory remit, address prioritized AI risks, calibrate rules based on regulatory experience and stakeholder input, and recommend whether the U.S. AI framework should prioritize other emerging risks (Table A4). The central risk function—ideally comprising scientific and technological experts, government officials, and independent private sector representatives—would evaluate the effectiveness of this framework, identify emerging AI risks, advise Congress and the federal government whether an intervention is required to address such risks, and if so, which regulators are best suited to address such emerging risks (Table A4).²⁷

While such proposals need to be more carefully evaluated and adjusted to suit the unique features of the U.S. regulatory architecture and policy objectives, they provide a useful starting point for thinking more strategically about ways to address future AI risks while maintaining a flexible regulatory approach. Furthermore, developing mechanisms to identify and address emerging AI risks would help improve public trust in AI and emerging technologies.

²⁵ DSIT, “Pro-Innovation Approach to AI.”

²⁶ Ibid.

²⁷ Nabil, “UK Approach to AI Governance.”

V. Strategies to Engage the Private Sector and Academic Institutions in AI Governance

The Trump administration should consider implementing mechanisms to engage the private sector and academic institutions more closely in AI governance. Such mechanisms are important for two reasons. First, the private sector and academic institutions have been instrumental in driving AI innovation. Second, given the rapidly evolving nature of AI-enabled technologies, the AI governance landscape is increasingly characterized by asymmetric information and a mismatch in technological expertise between regulators and the private sector. Developing mechanisms to continuously solicit feedback from external stakeholders when designing AI regulations is, therefore, crucial to maintaining a flexible regulatory approach.²⁸

Several policy tools could be incorporated into the U.S. AI framework to pursue closer engagement with private actors in developing AI rules. First, as discussed later, AI sandbox programs can help improve the regulatory understanding of emerging technologies and craft proportionate rules for AI applications in different sectors. Second, innovation hubs can be yet another source of information for startups and businesses to become aware of new commercial and investment opportunities, as well as compliance requirements associated with AI applications in different sectors.²⁹

Finally, soliciting feedback from businesses and monitoring the economic impact of AI regulations should also be part of the U.S. national AI strategy. To that end, AI working groups comprising regulators, academic and policy experts, and business representatives can provide an avenue for continued engagement between the private and public sectors in shaping AI governance.³⁰

VI. Well-Designed Artificial Intelligence Sandboxes to Improve the Regulatory Understanding of AI Technologies and Craft Flexible AI Rules

The U.S. government should consider developing multiple AI sandboxes to maximize the benefit of a flexible, innovation-focused approach to AI regulation. Such programs would allow companies to test innovative products and services under close regulatory supervision for a limited period while benefiting from regulatory waivers, expedited registration, and compliance guidance. Meanwhile, regulators would gain deeper insights into how emerging technologies and business models interact with existing laws and regulations. These insights would enable policymakers to craft more effective AI rules that foster technological innovation while mitigating risks.³¹

²⁸ Ryan Nabil, “Strategies to Improve the National Artificial Intelligence Research and Development Strategic Plan,” Competitive Enterprise Institute OnPoint, no. 282 (2022),

<https://cei.org/studies/strategies-to-improve-the-national-artificial-intelligence-research-and-development-strategic-pla/>

²⁹ Nabil, “How Regulatory Sandbox Programs Can Promote Innovation.”

³⁰ DSIT, “Pro-Innovation Approach to AI.”

³¹ Nabil, “How Regulatory Sandbox Programs Can Promote Innovation.”

Recognizing the innovation potential of AI sandboxes, several jurisdictions have introduced similar programs to craft a more flexible, innovation-friendly regulatory approach.³² Although the EU had initially expressed a lukewarm attitude towards regulatory sandboxes, it has since endorsed sandboxes—with the AI Act requiring each member state to establish (or join) at least one AI sandbox by August 2026.³³ The UK has been exploring various models for introducing AI sandboxes, while Singapore, Switzerland, and Norway have also launched similar initiatives.³⁴

To maximize their effectiveness, AI sandboxes must be carefully designed—a crucial consideration in the U.S. context, where regulatory fragmentation and the lack of coordination between federal and state regulators have hindered regulatory sandboxes in financial services. U.S. policymakers would benefit particularly from studying existing models for AI sandboxes with a view to develop and evaluate the regulatory designs of AI sandboxes programs that would best fit the U.S. regulatory context, as discussed in greater detail in the *Journal of Law, Economics, & Policy*. Based on this analysis, the U.S. government should consider establishing both multi-sector and sector-specific sandboxes to encourage AI innovation and calibrate AI rules accordingly. Finally, making AI sandboxes open to non-U.S. companies could help attract cutting-edge foreign startups and AI firms, further strengthening the U.S. position as a leading center in AI innovation.³⁵

VII. International AI Sandboxes to Promote Transatlantic Innovation and Cooperation

To maximize the benefits of AI sandbox programs, the United States should go one step further and design reciprocal AI sandboxes with like-minded countries such as France, Germany, Switzerland, and the UK. While no major jurisdictions have created such a program to the best of our knowledge, U.S. state legislation establishing state-level sandbox programs typically includes language indicating that state governments can create reciprocal sandbox arrangements with foreign regulators.³⁶ Reciprocal sandbox programs designed at the federal level would provide sandbox participants from signatory countries easier access to the equivalent U.S. regulatory sandboxes and vice versa.

Such programs could be particularly attractive to innovative foreign AI startups and companies that seek to understand and comply with U.S. regulatory requirements and enter U.S. markets. Likewise, reciprocal sandboxes could help U.S. businesses understand and comply with foreign regulatory frameworks, such as the EU's AI Act, and offer innovative products in those markets. By facilitating closer collaboration between foreign regulators and companies and facilitating harmonization of regulations and standards, reciprocal sandboxes could also help strengthen international economic and technology cooperation.

³² Laura Galindo-Romero, Karine Perset, and Francesca Sheeka, “An Overview of National AI Strategies and Policies,” Going Digital Toolkit Note, no. 14 (2021), https://goingdigital.oecd.org/data/notes/No14_ToolkitNote_AIStrategies.pdf.

³³ European Parliament and Council of the European Union, Regulation (EU) 2024/1689 of 13 June 2024 laying down harmonised rules on artificial intelligence, OJ L 327 (12 July 2024), art. 57.

³⁴ Ryan Nabil, “Artificial Intelligence Regulatory Sandboxes,” *Journal of Law, Economics & Policy* 19, no. 2 (2024): 295–348, <https://www.jlep.net/s/JLEP-192-Final.pdf>.

³⁵ Nabil, “Artificial Intelligence Regulatory Sandboxes.”

³⁶ Ibid.

VIII. Strengthened Bilateral Cooperation and Multilateral Engagement in AI Governance

Beyond AI sandboxes, the United States should consider other mechanisms—such as joint declarations, Executive agreements, and joint research programs—to strengthen tech cooperation at the bilateral level. In this context, the Joint U.S.-UK Declaration on Cooperation in AI Research and Development in September 2020 and the Atlantic Declaration in June 2023 were steps in the right direction.³⁷ Likewise, the U.S.-EU Digital Trade and Technology Council represents another forum through which the United States could pursue closer economic and technological cooperation with the EU and EU member states. Similar opportunities also exist for bilateral cooperation with Switzerland and Japan, which seek to adopt a flexible, light-touch approach to AI governance.³⁸ Establishing research partnerships—similar to Canada and the UK’s arrangements with Japan and the EU—could also help deepen U.S. technology cooperation with other advanced economies.

Ultimately, the United States needs to look beyond bilateral relationships and strengthen its multilateral engagement in global AI governance. Although the United States is part of several multilateral fora and institutions active in AI governance, such as the OECD and the Global Partnership on AI, the U.S. appears to punch below its weight in contributing to the development of international AI norms through these organizations. By participating more actively in such fora—as has been the case with Japan and the UK’s more multilateralist approach—the U.S. government can more actively contribute to the development of international AI norms and technical standards.³⁹

The development of such norms could be particularly beneficial for emerging-market and developing countries, many of which lack a robust AI governance infrastructure and look to international institutions to develop best practices in responsible AI. Along with like-minded partners—including the EU, the UK, Switzerland, Canada, and Japan—the United States could play a more active role in developing multi-stakeholder platforms for AI governance dialogues between state and private actors from both industrialized and emerging-market countries. As legislators and leaders in various jurisdictions seek to develop national AI strategies, the United States can be a leading voice for advocating a principles-based, innovation-focused AI approach that promotes economic growth and innovation while mitigating current and future AI risks.

³⁷ “The Atlantic Declaration: A framework for a twenty-first century US-UK Economic Partnership,” June 8, 2023, <https://www.gov.uk/government/publications/the-atlantic-declaration>. “Declaration of the United States of America and the United Kingdom of Great Britain and Northern Ireland on Cooperation in Artificial Intelligence Research and Development: A Shared Vision for Driving Technological Breakthroughs in Artificial Intelligence,” September 25, 2020, <https://www.gov.uk/government/publications/declaration-of-the-united-states-of-america-and-the-united-kingdom-of-great-britain-and-northern-ireland-on-cooperation-in-ai-research-and-development>.

³⁸ Staatssekretariat für Bildung, Forschung und Innovation [State Secretariat for Education, Research, and Innovation], “Herausforderungen der künstlichen Intelligenz: Bericht der interdepartementalen Arbeitsgruppe «Künstliche Intelligenz» an den Bundesrat” [“Challenges of Artificial Intelligence: Report of the Interdepartmental Working Group on Artificial Intelligence to the Federal Council”], December 2019, <https://www.sbfi.admin.ch/sbfi/de/home/bfi-politik/bfi-2021-2024/transversale-themen/digitalisierung-bfi/kuenstliche-intelligenz.html>. Der Bundesrat [The Federal Council], “Leitlinien «Künstliche Intelligenz» für den Bund: Orientierungsrahmen für den Umgang mit künstlicher Intelligenz in der Bundesverwaltung” [“Artificial Intelligence Guidelines for the Federal Government: Orientation Framework for Dealing with Artificial Intelligence in the Federal Administration”], November 2020, <https://www.sbfi.admin.ch/sbfi/de/home/bfi-politik/bfi-2021-2024/transversale-themen/digitalisierung-bfi/kuenstliche-intelligenz.html>. METI, “Governance Guidelines for Implementation of AI Principles.”

³⁹ METI, “Governance Guidelines for Implementation of AI Principles.”

Appendix

Table A1. Characteristics of the UK's Pro-Innovation AI Framework

Characteristic	Description
Pro-innovation	Enabling rather than stifling responsible innovation.
Proportionate	Avoiding unnecessary or disproportionate burdens for businesses and regulators.
Trustworthy	Addressing real risks and fostering public trust in AI in order to promote and encourage its uptake.
Adaptable	Enabling us to adapt quickly and effectively to keep pace with emergent opportunities and risks as AI technologies evolve.
Clear	Making it easy for actors in the AI life cycle, including businesses using AI, to know what the rules are, who they apply to, who enforces them, and how to comply with them.
Collaborative	Encouraging government, regulators, and industry to work together to facilitate innovation, build trust and ensure that the voice of the public is heard and considered.

Source: DSIT and UK Office for AI (2023).⁴⁰

Table A2. Guidelines for Regulators for Applying the UK's AI Framework

Scope	Description
Proportionate, context-specific, and flexible approach	Adopt a proportionate approach that promotes growth and innovation by focusing on the risks that AI poses in a particular context.
Prioritised risks and risk assessments	Consider proportionate measures to address prioritised risks, taking into account cross-cutting risk assessments undertaken by, or on behalf of, government.
Regulatory enforcement	Design, implement, and enforce appropriate regulatory requirements and, where possible, integrate delivery of the principles into existing monitoring, investigation, and enforcement processes.
Regulatory flexibility	Enabling us to adapt quickly and effectively to keep pace with emergent opportunities and risks as AI technologies evolve.
Awareness and transparency	Making it easy for actors in the AI life cycle, including businesses using AI, to know what the rules are, who they apply to, who enforces them, and how to comply with them.
Collaboration and public trust	Encouraging government, regulators, and industry to work together to facilitate innovation, build trust and ensure that the voice of the public is heard and considered.

Source: DSIT and UK Office for AI (2023).⁴¹

⁴⁰ DSIT, "Pro-Innovation Approach to AI."

⁴¹ Ibid.

Table A3. Categories of AI Risks Under the European Union's AI Act

Category and Requirement	Examples	Statutory Basis
Unacceptable risk: Prohibited	Social scoring, facial recognition, and dark-pattern AI	Art. 5
High risk: Conformity assessment	Education, employment, justice, immigration, and law	Art. 6 & ss.
Limited risk: Transparency	Chatbots, deep fakes, and emotional recognitions	Art. 52
Minimal risk: Code of conduct	Spam filters and video games	Art. 69

Source: Lilian Edwards, Ada Lovelace Institute (2022)⁴²

Table A4. Designing a U.S. Central Risk Function Mechanism for Artificial Intelligence Risks

Stakeholder	Identification*	Enforcement	Monitoring*
Congress and the Federal Government	i) Creates the AI framework to identify risks; ii) Decides which risks to tolerate, regulate, and prioritize.	Delegates the enforcement of the AI Framework to sectoral regulators.	Updates the statutory framework to address new risks if identified.
Central Risk Function Mechanism	i) Identify and prioritize new AI risks; ii) Provide recommendations if the new risks require government intervention.	i) Recommend which regulator(s) should address those risks; ii) Create overall risk assessment frameworks; iii) Provide advice to regulators on technical aspects of regulation; iv) Share AI regulatory best practices.	Monitors risks and reports them to Congress and the Executive.
Sectoral Regulators	i) Identify and prioritize sector-specific AI risks; ii) Evaluate whether newly identified risks should be prioritized and addressed.	i) Create regulatory guidance for businesses based on the central risk function's risk assessment framework; ii) Update regulatory guidelines and rules based on stakeholder feedback on how effectively they are working; iii) Take enforcement action for violations.	Reports on the effectiveness of addressing AI risks.
Businesses	Provide information to sectoral regulators and the central risk function, as necessary and appropriate.	Comply with regulatory guidance and rules and incorporate the risk assessment framework in internal practice.	Inform the relevant regulator(s) and the central risk function mechanism if risk mitigation measures fail to address the risks.

* The mechanisms highlighted in grey comprise a regulatory feedback loop between the federal government, sectoral regulators, the central risk function, and businesses subject to the AI framework to identify and mitigate emerging risks.

Source: Author based on DSIT and UK Office for AI (2023)⁴³

⁴² Edwards, "The EU AI Act."

⁴³ DSIT, "Pro-Innovation Approach to AI."