

July 7, 2023

Ms. Stacy Murphy
Deputy Chief Operations Officer/Security Officer
Office of Science and Technology Policy
Executive Office of the President
1650 Pennsylvania Ave, NW
Washington, DC 20504

Comments of Business Roundtable RE: “Request for Information: National Priorities for Artificial Intelligence”

Federal Register No. 2023–11346

Dear Ms. Murphy:

This letter is submitted on behalf of Business Roundtable, an association of more than 200 chief executive officers (CEOs) of America’s leading companies representing every sector of the U.S. economy. Business Roundtable CEOs lead U.S.-based companies that support one in four American jobs and almost a quarter of U.S. GDP. We appreciate the opportunity to respond to the Office of Science and Technology Policy’s (OSTP) Request for Information (RFI) on National Priorities for Artificial Intelligence (AI).

Introduction

Business Roundtable member companies across sectors—technology, communications, retail, financial services, health, public safety and security, defense, manufacturing, hospitality, insurance and others—rely on data and data-driven processes and solutions to create, deliver and improve innovative products and services across the United States and around the world. Rapid innovation and adoption of AI is transforming the nature of work across every industry and reshaping how people interact with and experience the world around them. AI technologies help businesses deliver smarter products and services and have enormous potential to drive broader positive change for Americans’ health, safety and prosperity.

Our members, who are among the world's largest developers and users of AI, recognize the critical importance of responsible AI practices. They are deeply committed to fostering the development and utilization of responsible AI systems that effectively mitigate risks, promote responsible innovation and engender trust from consumers, governments and the public to maximize the societal and economic benefits of AI while safeguarding public interest.

Business Roundtable applauds OSTP for soliciting comments on national priorities to advance our shared goals of U.S. competitiveness and leadership and we urge continued cooperation with the private sector as the process continues. As noted in our response to the National Telecommunications and Information Administration's Request for Comment on AI Accountability Policies, we believe that advancing safe and trustworthy AI is a shared responsibility among stakeholders, including between government and the private sector.¹ Below, we provide specific responses to OSTP's question categories posed in the RFI.

Responses to Question Categories Posed in OSTP's RFI

I. Protecting rights, safety, and national security

[Q1] What specific measures – such as standards, regulations, investments, and improved trust and safety practices – are needed to ensure that AI systems are designed, developed and deployed in a manner that protects people's rights and safety? Which specific entities should develop and implement these measures?

Unlocking the enormous potential of AI depends upon fostering public trust, understanding and widespread support for AI adoption — imperatives with shared responsibility across corporations and government. To that end, it is crucial that companies develop, deploy and use AI systems responsibly and with a careful eye towards promoting security, safety and human autonomy across applications and use cases.

In January 2022, Business Roundtable worked with member companies to launch its Responsible AI Initiative, publishing two foundational documents focused on both improved trust and safety practices for responsible AI and policy initiatives:

1. **Roadmap for Responsible AI²** (Roadmap), which sets out ten principles to guide companies of all sizes, in every sector and at every point along the AI value chain across three foundational characteristics of Responsible AI: (1) trusted and inclusive; (2) effective, safe and secure; and (3) accountable governance.
2. **Policy Recommendations³** for the U.S. government which encourages federal approaches to AI practices, rules and guidelines that build public trust in AI while enabling innovation and promoting continued U.S. leadership.

¹ See Business Roundtable, Response to National Telecommunications and Information Administration Request for Comment, NTIA-2023-0005 (June 12, 2023), <https://www.regulations.gov/comment/NTIA-2023-0005-1159>.

² See Business Roundtable, Roadmap for Responsible AI (Jan. 26, 2022), https://s3.amazonaws.com/brt.org/Business_Roundtable_Artificial_Intelligence_Roadmap_Jan2022_1.pdf

³ See Business Roundtable, Policy Recommendations for Responsible Artificial Intelligence (January 26, 2022), https://s3.amazonaws.com/brt.org/Business_Roundtable_Artificial_Intelligence_Policy_Recommendations_Jan2022_1.pdf

Specifically, Business Roundtable's Roadmap includes principles for organizations developing and deploying AI to ensure Americans' rights and safety are protected at every stage of the AI lifecycle, and that companies' AI oversight and governance will result in responsible AI deployment:

- Implement safeguards against unfair bias where AI systems may produce significant and high-consequence outcomes for individuals, while recognizing opportunities for AI to mitigate human bias.
- Explain the relationships between AI systems' inputs and outputs where possible and appropriate, particularly for systems that may result in significant and high-consequence outcomes for individuals, as well as the extent to which such inputs and outputs are governed by human oversight.
- Equip deployers of AI systems with sufficient information and training to support responsible and trustworthy downstream use.
- Disclose to end users when they are directly interacting with AI agents that simulate human interactions (e.g., chatbots).
- Evaluate and monitor model fitness and impact continuously to allow for adjustments for fitness for purpose, accuracy and resilience.

Business Roundtable members are already putting these principles into practice to create responsible AI systems and AI-enabled services by establishing cross-functional AI ethics and governance committees, conducting regular fairness and ethics assessments, and performance monitoring and refining existing internal risk management processes. For example, SAS' Viya platform⁴ supports a range of accountability mechanisms by providing clients with capabilities to assess data quality, detect and mitigate biases, assess model fairness, provide explainability, monitor system performance and support data privacy, all of which enables organizations to build employee and customer confidence in AI-supported systems.

Business Roundtable CEOs are committed to building trust in AI among our customers and stakeholders and to working toward the beneficial potential of AI for society and the economy by adopting and advocating for responsible AI governance. However, there is no single approach to designing, developing and deploying AI systems that protect individuals' rights and safety. For this reason, organizations should have the flexibility to design and implement specific measures that will most effectively protect individual rights and safety within their own specific use case contexts, consistent with clearly articulated frameworks and guidelines (e.g., NIST's AI Risk Management Framework) developed in partnership with the public sector and with existing consumer protection laws and regulations.

⁴ *Ibid.*

To the extent that measures are implemented through government rules, regulations or standards, it is important to remember that the context and corresponding risk levels of AI applications exist in a wide spectrum. As such, requirements should be outcome-focused and take a risk-based approach to avoid over-regulating uses of AI which have no significant impact on individuals or do not pose potential for societal harm. For cases in which government action is necessary to protect people's rights and security, such action should be consistent and compatible with the following principles:

- Any regulatory approaches to AI considered or adopted in the United States should be contextual, risk-based, proportional and use-case specific. Any frameworks, guidance and regulation should be tailored to specific AI use cases, rather than broadly regulating any technology or application outright and should be appropriately calibrated depending upon the risk of substantial harm.
- AI measures should incentivize good-faith and demonstrated efforts to adhere to requirements, norms and standards.
- In developing these measures, policymakers should conduct a thorough assessment of existing regulatory gaps before establishing new regulations to avoid overlapping or inconsistent rules and to understand where guidance is most needed.
- AI measures should include clear definitions with case study examples informed by continued dialogue with industry stakeholders.
- Policymakers should explore the use of evidence-based regulatory approaches and tools that allow for the iteration of governance practices (e.g., regulatory sandboxes) and opportunities for industry to discover and share best practices.
- Finally, government could incentivize industry to engage in self-assessments, whether such work is performed internally or against external guidelines or standards.

Given the important role of data to the development of AI and to realizing the benefits of AI, data privacy is critical for responsible, safe AI. U.S. data privacy laws are increasingly fragmented across industries, geographies and jurisdictions, creating confusion among consumers and a complicated web of compliance activities for companies, thereby further detracting already limited time and resources from strategic and compliance work on AI governance efforts. Business Roundtable strongly supports a national consumer privacy law, which would strengthen protections for consumers across the country, while offering Congress the opportunity to create a holistic, preemptive approach to privacy and data security.

[Q2] How can the principles and practices for identifying and mitigating risks from AI, as outlined in the Blueprint for an AI Bill of Rights and the AI Risk Management Framework, be leveraged most effectively to tackle harms posed by the development and use of specific types of AI systems, such as large language models?

Business Roundtable supports the shared principles of innovation, transparency, safety, trustworthiness and inclusion, and strongly believes in the importance of continued public-

private engagement as the Administration develops principles and practices for detecting and mitigating risks from AI.

The NIST AI Risk Management Framework (RMF) is a strong example of a successful, collaborative process that included a diverse range of stakeholders from research institutions, AI developers and users, and the broader technology industry. The core values of the RMF — voluntary, risk-based and rights-preserving — are very aligned with the principles set forth in Business Roundtable’s Roadmap. It is important that other frameworks and guidance developed by the U.S. government are consistent with the RMF to avoid uncertainty and fragmentation across the federal approach to AI. As noted in the RMF, and consistent with the Roadmap, AI accountability frameworks should adapt to the AI landscape as technologies continue to develop.

The resulting work provides AI developers and users with a risk-based guide to incorporating transparency and accountability throughout the entire AI lifecycle. Leaning on the guidance in the NIST AI RMF, policymakers should consider the importance of a proportionate, risk-based approach, with higher-risk applications subject to more in-depth reviews, robust cost-sensitive analyses and stricter (often cross-domain) governance mechanisms (e.g., ethics committees). Importantly, the RMF is an evolving document, and future iterations can account for the best and most timely guidance on implementing accountability efforts as AI technology advances. This could include ongoing collaboration between NIST and the private sector to (a) develop new use case profiles to account for new adaptation of large language models (similar to the other use case profiles), and (b) add specificity and clarity regarding operationalization of the socio-technical aspects of the RMF, which would help companies calibrate their compliance efforts and encourage innovative approaches to bias mitigation.

Business Roundtable companies are at the forefront of putting the NIST AI RMF principles and processes into practice. Many companies are embedding responsible AI principles into their internal governance structures, and some businesses are helping their clients to achieve AI systems that align with the principles laid out in the Roadmap:

- Booz Allen’s aiSSEMBLE™ – a lean manufacturing approach to AI engineering designed to simplify the engineering and deployment of AI systems – embeds RAI principles and practices (e.g., transparency, traceability, auditability, dynamic data and model drift detection) into AI system design, execution and monitoring—ensuring that RAI is built-in from the start. Booz Allen is now working to integrate these capabilities with their recently announced investment in Credo.ai to further extend their RAI solution offerings in market.
- IBM’s Ethics by Design (EbD) Framework integrates tech ethics into the organization’s full tech development pipeline, including AI, and embeds responsible governance across the organization.

[Q4 and Q7] What are the national security benefits associated with AI? What can be done to maximize those benefits? What are the national security risks associated with AI? What can be done to mitigate these risks?

AI is transforming and revolutionizing businesses across all industries with increasingly widespread use as the technology advances and becomes more accessible. The continued innovation and growth in AI technology will be critical to ensuring that the United States is able to leverage it to protect against attacks by malicious actors, who may also choose to leverage it regardless of legal limitations. Some examples of how AI can advance national security include:

- **Cyber Defense and Influence:** AI can significantly bolster cyber defenses and intelligence processing by increasing the speed of incident response; detecting and preventing threats; sifting through massive amounts of data to spot patterns and elevate the most important alerts for human attention; helping to identify the root cause of attacks; etc.
- **Autonomy and Intelligence:** In addition to the automation of simple tasks, AI capabilities in robotics, computer vision and natural language processing can be used in consequential settings like intelligence gathering and enhancing the safety and reliability of critical systems and equipment.⁵

While AI is an increasingly essential component of cyber and data security defenses, its use also carries some national security risks associated with AI, which experts are still working to fully grasp. AI systems – like any technology – can be vulnerable to adversarial attacks. AI capabilities also may be used by malicious actors to launch sophisticated cyber-attacks. In addition, as AI applications often require access to and processing of vast amounts of data, AI can create privacy and surveillance risks.

Accessible and appropriately tailored training for AI developers and users can help to mitigate these risks, including education about social engineering and understanding cyber threats. In addition, to protect user data and the national interest, organizations should explore the use and efficacy of privacy enhancing technologies in AI contexts, and governments should encourage strong data protection frameworks to safeguard individual privacy rights.

To both fully understand the national security risks of AI and to better manage them, the U.S. government, in partnership with the private sector, should continue to engage in collaboration and the establishment of common standards with value-aligned countries and in international standards organizations. Governments, research institutions and industry stakeholders should

⁵ See Greg Allen Taniel Chan, Belfer Center for Science and International Affairs, *Artificial Intelligence and National Security* (July 2017), <https://www.belfercenter.org/publication/artificial-intelligence-and-national-security#:~:text=Advances%20in%20AI%20will%20affect,a%20broader%20range%20of%20actors>

continue to collaborate to develop global frameworks for responsible AI use, fostering transparency, interoperability, information sharing and coordinated responses to security risks, and to understand which of these fora are best positioned to lead coordination efforts.

Maximizing the national security benefits and minimizing the risks associated with AI will require a multi-faceted approach involving technological advancements, regulatory frameworks, international cooperation and responsible practices across all stakeholders.

II. Advancing equity and strengthening civil rights

[Q10] What are the unique considerations for understanding the impacts of AI systems on underserved communities and particular groups, such as minors and people with disabilities? Are there additional considerations and safeguards that are important for preventing barriers to using these systems and protecting the rights and safety of these groups?

Mitigating unfair bias and harm is a core component of Business Roundtable's Roadmap. The three foundational tenets that the Roadmap is built on support this aim: first, a respect for safety, dignity and autonomy of customers, end users and employees; second, a belief that concepts of effectiveness, fairness and security are inherent to Responsible AI; and third, an understanding that humans play a vital role across the AI lifecycle. Ensuring alignment with Responsible AI principles requires AI systems to be both trusted and inclusive.

Organizations should strive to assemble AI system design, development and deployment teams that represent a diversity of professional experience, subject matter expertise and lived experience — consistent with broader diversity goals and efforts. Similar diversity objectives should extend to ethics committees, governance boards or any internal authority that oversees AI. Organizations should also consider how and when diverse perspectives can supplement internal knowledge, particularly for high-consequence systems with direct human impact.

AI systems should also be designed and implemented in a manner that promotes transparency, explainability and interpretability. The ability to understand the relationship between AI inputs and outputs, for both AI deployers and end users, serves as a check against unfair bias. Both AI workers and end users should be equipped with the tools and knowledge to responsibly interact with AI. End users should be able to know when they are interacting with AI systems and understand outputs from this system — particularly for systems with high-consequence outcomes for individuals.

Additionally, comprehensively addressing the impacts of AI systems on underserved communities and groups will require private-public sector collaboration to reskill and upskill employees, support education and skill-building to accelerate the development of a modern and diverse workforce, and make educational resources and training programs more accessible across diverse geographic, demographic and socio-economic backgrounds.

Finally, and most importantly, the individual and collective project of building and maintaining trust through responsible AI is a continuous and evolving journey. Collecting feedback, monitoring and evaluating system performance, and integrating emerging best practices can ensure that AI systems and applications are working as intended and generating benefits for all groups and communities.

[Q12] What additional considerations or measures are needed to assure that AI mitigates algorithmic discrimination, advances equal opportunity and promotes positive outcomes for all, especially when developed and used in specific domains (e.g., in health and human services, in hiring and employment practices, in transportation)?

Business Roundtable believes that responsible design, development, deployment and use of AI is critical to building consumer, government and public trust while advancing U.S. leadership in these emerging technologies and using them to the benefit of society and the economy. Our companies are at the forefront of these issues and are leading efforts to create and implement responsible AI governance, risk management, accountability and transparency. Organizations should implement safeguards against unfair bias where AI systems may result in significant and high-consequence outcomes for individuals. Data collection and use is a core feature of AI systems, and likewise requires responsible and accountable processes to avoid introducing algorithmic bias. In certain areas where best practices and established conventions are not clear (e.g., solutions and frameworks for mitigating proxy bias), additional government guidance can help to clarify effective processes for managing bias risks and is a necessary precursor to any imposition of requirements.

III. Bolstering democracy and civic participation

[Q16] What steps can the United States take to ensure that all individuals are equipped to interact with AI systems in their professional, personal and civic lives?

Recommendations on AI education, training and awareness are a key component of Business Roundtable's Policy Recommendations for AI. The U.S. government should partner with industry to build AI literacy and relevant skill sets across the country to invest in AI education and proficiency at all levels. This includes supporting AI education at academic and trade institutions to broaden AI knowledge and prepare students for AI-compatible roles; developing early education curricula and consumer literacy programs; and making AI educational resources and training programs widely accessible across geographic areas and socio-economic backgrounds.

Business Roundtable also supports public-private collaboration on industry training and reskilling efforts, including multisectoral partnerships among education institutions, industry and government entities to promote applied AI learning and apprenticeships as well as investments in talent recruitment to enhance technical AI capacity across federal agencies.

IV. Promoting economic growth and good jobs

[Q17] What will the principal benefits of AI be for the people of the United States? How can the United States best capture the benefits of AI across the economy, in domains such as education, health and transportation? How can AI be harnessed to improve consumer access to and reduce costs associated with products and services? How can AI be used to increase competition and lower barriers to entry across the economy?

AI is an incredibly versatile technology with wide-ranging applications that will have far-reaching and impactful effects, transforming the way individuals interact and organizations operate to enhance health, safety and productivity. AI can also deliver tangible societal benefits, such as improving government services, increasing accessibility for individuals with disabilities and reducing unconscious bias to drive more equitable outcomes.

AI is critical to fraud detection and prevention to protect Americans and their pocketbooks, especially given the growing sophistication of attacks. For example, Visa's Advanced Authorization (VAA)⁶ combines Visa's proprietary online model with offline neural network-based machine learning to evaluate fraud risk across its entire network (VisaNet) in real-time, helping to mitigate potential fraud risk before transactions go through.

In many cases, AI systems improve existing alternatives by reducing unconscious biases, improving safety and driving more equitable outcomes for consumers and communities. These opportunities should be accompanied by effective AI accountability mechanisms to address relevant risks of harms and expand access to important services.

[Q18 and Q20] How can the United States harness AI to improve the productivity and capabilities of American workers, while mitigating harmful impacts on workers? How can the United States promote quality of jobs, protect workers and prepare for labor market disruptions that might arise from the broader deployment of AI in the economy?

Business Roundtable supports public and private sector investments in a future-ready workforce, including the need to upskill, reskill and expand opportunities. Proactive investments in a dynamic workforce, including public-private partnerships, are crucial to ensuring that the United States and its workforce can realize the full benefits of AI. Our Roadmap highlights these investments and makes specific recommendations for organizations and government to bring new opportunities and create new jobs, including:

- Determine which tasks AI will augment, change or eliminate, and consider where new tasks and jobs may be created;

⁶ See Business Roundtable, Business Roundtable Companies Put AI Recommendations into Practice (Jan. 25, 2023), <https://www.businessroundtable.org/ai-innovation-at-work-putting-principles-into-practice>.

- Support employees whose roles and responsibilities may change, and plan to reskill, upskill and/or provide new opportunities;
- Support education and skill-building efforts to accelerate the development of a modern and diverse workforce capable of developing and using AI responsibly; and
- Make educational resources and training programs accessible across geographic, demographic and socio-economic backgrounds to broaden the AI talent pipeline.

Dell Technologies, for example,⁷ works with site managers to identify and diagnose specific points along the production process where there is potential to optimize human performance by integrating machine learning models. This AI-enabled process does not eliminate the need for humans in the manufacturing process, but rather frees them up to do what they do best: conduct careful, pinpoint analysis of specific defects, while leaving the high-frequency work of flagging defects to the AI system. The result is that manufacturers can achieve higher quality, more consistency and increased productivity without sacrificing job satisfaction or increasing the stress and strain on their workforce.

V. Innovating in public services

[Q26] How can the Federal Government work with the private sector to ensure that procured AI systems include protections to safeguard people's rights and safety?

The U.S. government can and should draw on best practices for AI governance developed by industry, including Business Roundtable's Roadmap for Responsible AI, alongside well-established risk assessment frameworks such as the NIST AI RMF. Lack of clear guidelines from the federal government regarding agencies' use of AI prevents smooth procurement because of uncertainty regarding governance. Uniform or recommended assessment frameworks for AI governance could help public servants more quickly spot and address ethical concerns that could arise throughout the acquisition process, as well as train them to incorporate this thinking from the outset when approaching future projects. Industry and government could also partner to train public servants on enhancing awareness of AI's opportunities and risks, and how to implement AI in practical terms.

Conclusion

Business Roundtable looks forward to continued engagement with OSTP and other thought leaders and policymakers on these important topics. To discuss our response or these issues at any time, please contact Amy Shuart, Vice President of Technology & Innovation, Business Roundtable, at ashuart@brt.org or 202-496-3290.

⁷ *Ibid.*