

RECENTERING PUBLIC VALUES IN AI GOVERNANCE: EXAMPLES FROM THE BIDEN ADMINISTRATION

Deirdre K. Mulligan[†] & Kenneth A. Bamberger^{††}

ABSTRACT

This Article situates key Biden-Harris Administration AI initiatives within a “governance-by-design” framework—an approach we previously developed that centers public values, sectoral expertise, and participatory policymaking in decisions to regulate through technology. Governance-by-design argues for reorienting AI governance around three core principles: (1) privileging human and public rights by empowering domain-specific agencies while building a shared set of tools and approaches for risk assessment; (2) expanding agencies’ technical expertise through public hiring and multisector collaboration; and (3) preserving the publicness of policymaking through designs that foreground embedded values and embedding stakeholder engagement and impact evaluation throughout AI system development and deployment.

The Article uses three examples to illustrate how key Biden-Harris Administration AI actions reflect these governance-by-design principles:

- the Administration’s layered regulatory strategy that empowers sectoral agencies to safeguard human rights and public safety in AI use;
- the expansion of AI and rights-based expertise within government and the establishment of collaborative structures for risk management and evaluation; and,
- the institutionalization of practices that surface and interrogate the normative assumptions embedded in AI systems, while scaffolding public participation throughout their lifecycle.

We argue that together these initiatives offer an alternative to prevailing AI governance debates—particularly the dichotomy between risk-based and rights-based approaches, and the call for a centralized AI regulator. Instead, such governance-by-design provides a field-centric model that leverages existing institutional capacities, protects democratic norms, and re-centers the public in the often-private domain of AI development. It offers a durable, epistemically responsible framework for regulating AI systems in a way that supports both human rights and legitimate democratic governance.

TABLE OF CONTENTS

DOI: <https://doi.org/10.15779/Z38416T25K>

© 2025 Deirdre K. Mulligan and Kenneth A. Bamberger.

[†] Professor, UC Berkeley School of Information; Co-Faculty Director, Berkeley Center for Law & Technology; former Principal Deputy U.S. Chief Technology Officer at the White House Office of Science and Technology Policy, and Director of the National Artificial Intelligence Initiative Office (NAIIO), 2023–2024.

^{††} The Rosalinde and Arthur Gilbert Foundation Professor of Law, UC Berkeley; Co-Faculty Director, Berkeley Center for Law & Technology.

Much gratitude to Rachel K. Mucha for her superb research assistance.

I.	INTRODUCTION	1136
II.	THE MISDIRECTION OF AI GOVERNANCE DEBATES.....	1139
A.	THE RISKS VS. RIGHTS BINARY: HOW TO GOVERN.....	1140
B.	INSTITUTIONAL DESIGN IN AI REGULATION: WHO SHOULD GOVERN?.....	1143
III.	OUR FRAMEWORK FOR RECENTERING PUBLIC VALUES IN TECHNOLOGY GOVERNANCE: AN ALTERNATIVE TO THE AI DEBATES	1148
IV.	REFLECTING OUR GOVERNANCE PRINCIPLES: EXAMPLES FROM THE BIDEN-HARRIS ADMINISTRATION'S APPROACH TO AI GOVERNANCE.....	1151
A.	PRIVILEGING HUMAN AND PUBLIC RIGHTS: MAINTAINING EXPERT AGENCY AUTHORITY WHILE BUILDING A SHARED KNOWLEDGE BASE FOR RISK ASSESSMENT METHODS AND PRACTICES.....	1153
B.	BRINGING EXPERTISE AND CAPACITY INTO GOVERNMENT: DIRECT HIRING AND STAKEHOLDER INVOLVEMENT.....	1162
1.	<i>Bringing AI and AI Enabling Talent into Federal Service.....</i>	1162
2.	<i>Building the Responsible AI Field.....</i>	1166
C.	MAINTAINING THE PUBLICNESS OF POLICYMAKING: FOCUSING ON IMPACT RATHER THAN SYSTEMS AND REQUIRING STAKEHOLDER PARTICIPATION THROUGHOUT THE AI LIFECYCLE.....	1171
1.	<i>Reframing The Project of AI Governance</i>	1171
a)	The AI Bill of Rights	1171
b)	OMB Guidance to Federal Agencies	1175
2.	<i>The National Telecommunications and Information Administration (NTIA) and Model Weights</i>	1180
V.	CONCLUSION	1183

I. INTRODUCTION

Artificial Intelligence (AI) design and deployment displays attributes of a type that persistently confounds public governance. AI is not amenable to traditional command-and-control regulation reliant on uniform ex ante rules requiring certain conduct. Technical expertise resides largely in private rather than public actors and institutions. Engineers in the private sector make granular decisions regarding AI design. Private firms often manage the deployment of AI even when used by governments. Model developers and deployers, as well as workers and the public who use or whose rights and interests are affected by AI, possess the information needed to understand the

ways in which AI affects different segments of our society. And the nature of AI itself compounds the challenges of opacity, comprehension, and uncertainty that threaten to turn “over key policy questions to privately developed algorithmic systems.”¹

How, then, can we build a regulatory system that doesn’t outsource policy decisions? One that is equipped to meaningfully protect rights and safety. One that withstands the corrosive and insidious way private sector processes can dilute or undermine public goals to fit within management practices. A regulatory system in which technological opacity doesn’t prevent agency experts from applying their knowledge or preclude the public from participation.

In sum, how can we recenter public values in AI governance?

Elsewhere, we have taken a hard look at efforts to enlist technology to protect values in information and communication technology.² We concluded that “governance-by-design”—the purposeful effort to use technology to embed values or policies—had become a central mode of policymaking, but also that our existing regulatory system was fundamentally ill-equipped to prevent that phenomenon from subverting public governance.³

Specifically, we provided examples that showed how governance-by-design had undermined important governance norms and chipped away at our voting, speech, privacy, and equality rights.⁴ We further described the structural limitations of traditional legal and governance bodies that contributed to this problem. These include the limited technical expertise of many policy making bodies; the absence of a venue for policymakers to have the meta-discussion about when and whether it is appropriate to enlist technology in the service of values at all; and, relatedly, if technology is to be so used, how to prioritize among values.⁵

These structural limitations, we explained, contributed to processes that subvert fundamental democratic norms of intentional, deliberative,

1. Deirdre K. Mulligan & Kenneth A. Bamberger, *Procurement as Policy: Administrative Process for Machine Learning*, 34 BERKELEY TECH. L.J. 781, 807 (2019) [hereinafter Mulligan & Bamberger, *Procurement as Policy*].

2. See generally Deirdre K. Mulligan & Kenneth A. Bamberger, *Saving Governance-By-Design*, 106 CALIF. L. REV. 697 (2018) [hereinafter Mulligan & Bamberger, *Governance-By-Design*]; Mulligan & Bamberger, *Procurement as Policy*, *supra* note 1.

3. Mulligan & Bamberger, *Governance-By-Design*, *supra* note 2, at 697; see also Kenneth A. Bamberger, *Technologies of Compliance: Risk and Regulation in a Digital Age*, 88 TEX. L. REV. 669, 675–76 (2010) (describing the ways that “the use of technology systems to hardwire compliance” can “raise what might be called administrative-law concerns—concerns regarding the subversion of public norms requiring transparency, public oversight, and accountability in the exercise of regulatory discretion.”).

4. Mulligan & Bamberger, *Governance-By-Design*, *supra* note 2, at 698.

5. *Id.* at 750.

participatory, and expert public decision-making, free from capture or caprice. And they produced overbroad “technological fixes” that privilege singular values and often disfavor human rights.⁶ We further argued that the use of technology to regulate without addressing the limitations of existing legal policy and enforcement processes. This allowed governments and private actors to mask their aims, led to both unwitting and intentional privileging of some values over others, and obscured the resulting policy outcomes from ongoing public scrutiny as they recede into the technical desiderata of the technical environment.⁷

Finally, we proposed a framework for “saving” governance-by-design that emphasized a set of approaches that together could align technology governance with the norms of public governance and therefore be more responsive to public values, and not just private interests.⁸

Three are regulatory principles:

- (1) Privileging Human and Public Rights;
- (2) Ensuring that Regulators Possess the Right Attributes, including Broad Authority and Competence, as well as technical expertise; and
- (3) Maintaining the Publicness of Policymaking.⁹

The fourth is a design principle in service of the others, counseling modesty and restraint in design that wherever possible preserves flexibility rather than fixing values.¹⁰

This Article places that framework within the context of current debates regarding the appropriate metrics and institutional structure for AI governance: specifically (1) disputes over whether AI regulation should be “risk-based” or “rights-based,” and (2) arguments that a new, dedicated regulatory body should be created to administer the governance of AI. It identifies the ways those arguments are misdirected and explains the ways that our technology and design governance principles offer an alternative emphasis. With these three regulatory principles in mind, this Article considers and frames examples of AI governance approaches taken by the Biden-Harris Administration. Those examples follow regulatory principles that seek to ensure AI is developed and implemented in ways that support both democratic values and human rights, as well as the public’s safety and security: that the “public” is recentered in the often-private endeavor of AI development and

6. *Id.* at 739.

7. *Id.* at 721.

8. *Id.* at 705.

9. *Id.*

10. *Id.*

implementation. They reflect a field-centric, in contrast to what one of us has called “model-centric,” model of AI governance. This approach centers and maintains the meaning-making processes, rules, and norms that guide interactions and decisions within fields and protects them against the epistemological and other displacements that too often are a byproduct of automating or informing¹¹ activities. They thereby “support epistemically responsible behaviour.”¹² The approach addresses the inherent shortcomings of government technical expertise while at the same time reasserting agency domain expertise, and the rights and logics that undergird legitimate public governance processes.

These Biden-Harris Administration initiatives, then, concretize a suite of regulatory approaches for successfully recentering public values in AI governance. Such public values include (1) privileging human and public rights by maintaining expert agency authority while building a shared knowledge base for risk assessment methods and practices—a method for appreciating risk, while identifying which rights we seek to privilege; (2) bringing expertise and capacity into government through direct hiring and stakeholder involvement; and (3) maintaining the publicness of policymaking by evaluating impacts not technical systems and requiring stakeholder participation throughout the process.

II. THE MISDIRECTION OF AI GOVERNANCE DEBATES

Much of the current discourse around AI governance involves two debates. The first involves a binary discourse over the proper mode of AI regulation: specifically, whether AI should either be regulated in a manner that is focused on risks or, by contrast, one that places attention on rights. The second debate engages the question of governance institutions: whether addressing the risks of AI would be best accomplished through a new agency with technical expertise dedicated to regulating AI models or systems.¹³

11. SHOSHANA ZUBOFF, *IN THE AGE OF THE SMART MACHINE: THE FUTURE OF WORK AND POWER* (Basic Books 1988).

12. Judith Simon, *Distributed Epistemic Responsibility in a Hyperconnected Era*, in *THE ONLIFE MANIFESTO: BEING HUMAN IN A HYPERCONNECTED ERA* 145, 155–58 (Luciano Floridi ed., 2015).

13. An examination of the various proposals to create new U.S. agencies to regulate digital markets is outside the scope of this paper. For examples of such proposals, see Harold Feld, *The Case for the Digital Platform Act: Market Structure and Regulation of Digital Platforms*, ROOSEVELT INST. 17 (May 2019) (urging the United States to “either empower an existing agency or create a new agency to use these powers as necessary”); Tom Wheeler, Phil Verveer & Gene Kimmelman, *New Digital Realities; New Oversight Solutions in the U.S.: The Case for a Digital Platform Agency and a New Approach to Regulatory Oversight*, SHORENSTEIN CTR. 2 (Aug. 2020) (advocating for the creation of “a new Digital Platform Agency” to take an “agile approach to oversight built on risk management” that would include “cooperatively developed and

A. THE RISKS VS. RIGHTS BINARY: HOW TO GOVERN

Margot Kaminski has documented a “growing convergence” around the use of risk-based frameworks for AI governance.¹⁴ Regulators, policymakers, and scholars have identified this suite of approaches as particularly appropriate in regulatory contexts in which the implementation of policy goals is “technically and legally opaque.”¹⁵ Requiring assessments of systemic risk level by context promises rigor in allocating regulatory focus,¹⁶ as at least “[i]n its idealized form, risk-based regulation offers an evidence-based means of targeting the use of resources and of prioritizing attention to the highest risks in accordance with a transparent, systematic, and defensible framework.”¹⁷ Discerning levels of risk by context, moreover, allows regulators to choose to take on different levels of risk in light of broader benefits to society.¹⁸ The EU AI Act’s call for a “risk-based approach” in order to effect a “proportionate” set of binding rules for AI systems, for example,¹⁹ reflects its “underlying objective” “to strike an optimal (or proportionate) balance between innovation and the benefits of AI systems on the one hand, and the protection of fundamental values such as safety, health and fundamental rights on the

enforceable code of conduct for specific digital activities”), https://shorensteincenter.org/wp-content/uploads/2020/08/New-Digital-Realities_August-2020.pdf; Digital Platform Commission Act of 2023, S. 1671, 118th Cong. (2023) (creating an expert federal agency to comprehensively regulate digital platforms to protect consumers, promote competition, and defend the public interest). For an overview of calls for the creation of new agencies to regulate privacy and a wide range of other platform and AI-related issues (search, robotics, algorithms, etc.), see Asad Ramzanali, *Toward a Privacy Agency: Policy and Politics* Appendix V (Apr. 6, 2021) (M.P.P. Policy Analysis, Harvard Kennedy School of Government), <https://www.dropbox.com/scl/fi/zfw6bavrlxsd2r3vnf0gi/Toward-a-Privacy-Agency-Policy-and-Politics-PAE-Asad-Ramzanali-4.6.21.pdf?rlkey=ihibb9qpkcd0b5u8g7cfl8ty9&e=1&st=45cq8jmt&dl=0>.

14. Margot E. Kaminski, *Regulating the Risks of AI*, 103 B.U. L. REV. 1347, 1347 (2023).

15. *Id.* at 1365–66.

16. See *EU AI Act: First Regulation on Artificial Intelligence*, EUROPEAN PARLIAMENT, <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence#:~:text=AI%20regulation%20in%20Europe%3A%20the%20first%20comprehensive%20framework,-In%20April%202021&text=AI%20systems%20that%20can%20be,or%20less%20AI%20compliance%20requirements> (last updated Feb. 19, 2025) (“In April 2021, the European Commission proposed the first EU artificial intelligence law, establishing a risk-based AI classification system. AI systems that can be used in different applications are analysed and classified according to the risk they pose to users. The different risk levels mean more or less AI compliance requirements.”).

17. Julia Black & Robert Baldwin, *Really Responsive Risk-Based Regulation*, 32 L. & POL’Y 181, 181 (2010).

18. See generally Kaminski, *Regulating the Risks of AI*, *supra* note 14.

19. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 12 July 2024 on artificial intelligence, recital 26, O.J. (L 1689).

other.”²⁰ In a regulatory arena in which the prescription of detailed mandates is unfeasible, moreover, risk-based regulation’s focus on outcomes and assessments leaves flexibility in how goals are met. This approach thus often allows organizations to tailor their compliance measures to specific risks and contexts, enlisting them to document their process decisions and success or failure at doing so.

By contrast, rights-based approaches foreground concerns related to fundamental rights and individual fairness. Rights-based AI governance frameworks treat rights such as nondiscrimination and privacy as fundamental and inviolable.²¹ Legal scholars and policymakers in the rights-based camp point favorably to regulatory frameworks like the EU’s GDPR, which follows a “binary logic.” They provide a “minimum and non-negotiable level of protection” against certain harms of AI for all individuals, and a regulated entity’s action either provides this protection or fails to do so.²² Rights-based frameworks apply the same rules to everyone irrespective of the level of risk or harm. In the context of the GDPR, for example, a data processing action either provides users adequate protection from risk or harm or it falls short—there is no balancing of interests or level of harm that is tolerable at the individual or systemic level.²³

We and other scholars of regulation have cautioned against the excesses of both binary approaches to technology governance. While risk-based regulation is valuable in contexts featuring easily measured harms, reliance on a “comprehensive, defined, ex ante, body of regulatory mandates”²⁴ to govern technology design can leave unanswered the question of *risk to what values* in contexts involving “big, often-unquantifiable, often-contested, often-contextual, and often-individualized ‘risks.’”²⁵ As law-and-technology scholars Julie Cohen and Ari Waldman explain, regulatory oversight premised on an ex ante focus on risk mitigation can lead to a form of “regulatory managerialism” that embraces a narrow toolkit of risk modeling, digital control systems, and data analytics that ultimately focuses narrowly on efficiency and process values.²⁶ Regulating private activity around technology with these operational

20. Martin Ebers, *Truly Risk-Based Regulation of Artificial Intelligence: How to Implement the EU’s AI Act*, 16 EUR. J. RISK REGUL. 684, 685 (2025).

21. *Fundamentals of a Human Rights-Based Approach to Generative AI*, BSR (Feb. 2025), <https://www.bsr.org/files/BSR-Fundamentals-of-a-Human-Rights-Based-Approach-to-Generative-AI.pdf>.

22. RAPHAËL GELLERT, *THE RISK-BASED APPROACH TO DATA PROTECTION 2* (2020).

23. *Id.* An alternate take on the GDPR is that it is focused on a set of actions that were a priori determined to be high-risk.

24. Mulligan & Bamberger, *Governance-By-Design*, *supra* note 2, at 739.

25. Kaminski, *Regulating the Risks of AI*, *supra* note 14, at 1378–79.

26. Julie E. Cohen & Ari Ezra Waldman, *Introduction: Framing Regulatory Managerialism as an Object of Study and Strategic Displacement*, 86 L. & CONTEMP. PROBS. i, ix-x (2023); *see generally*

values largely in mind, we have explained elsewhere, can lead to a “hands-off deference” to values choices²⁷ and a “skewing of public legal norms by private interests.”²⁸ Risk assessment, “a key technique of managerialism directed at formalizing and constraining agency decision-making,”²⁹ was developed as “a political technology intended to discipline agencies, rather than a tool for revealing truths about the world.”³⁰ Risk-based regimes, accordingly, facilitate a “light-touch”³¹ approach to governance, while the centering of public values demands instead “activist” and “dynamic” regulators.³²

In the end, managerial approaches to technological “governance in private hands can produce symbolic or ceremonial structures that imbue corporate acts with apparent legitimacy but do little to further the public values at stake.”³³ In this way, as Kaminski describes, a risk-based metric for AI governance alone fails to account for “dignitary and justificatory concerns about algorithmic decision-making.”³⁴ A legitimate AI governance regime must account in a thick manner for the ways in which it identifies and protects the rights and public values that are at risk.

At the same time, while we have advocated for the importance of privileging human and public rights in technology governance, we have also raised cautions about the method for doing so. In particular, we have pointed to the “range of human rights and other public values”³⁵—some of which might be in tension—that could be embedded in technology design. The choice between them, we have argued, must derive from democratically- and

Frank A. Pasquale, *Power and Knowledge in Policy Evaluation: From Managing Budgets to Analyzing Scenarios*, 86 L. & CONTEMP. PROBS. 39, 43 n.20 (2023) (“There is, as Robert Post has observed, a critical distinction between governance and management.” (citing Robert Post, *Between Governance and Management: The History and Theory of the Public Forum*, 34 UCLA L. REV. 1713, 1788 (1986))).

27. Bamberger, *Technologies of Compliance*, *supra* note 3, at 684.

28. *Id.* at 726.

29. William Boyd, *With Regard for Persons*, 86 L. & CONTEMP. PROBS. 101, 104 (2023).

30. *Id.*

31. Cohen & Waldman, *supra* note 26, at xv.

32. Bamberger, *Technologies of Compliance*, *supra* note 3, at 735; *see also* KENNETH A. BAMBERGER & DEIRDRE K. MULLIGAN, *PRIVACY ON THE GROUND* 187, 225 (2015) (documenting the importance of such “activist” regulators in outcomes-based governance regimes).

33. Deirdre K. Mulligan & Kenneth A. Bamberger, *Allocating Responsibility in Content Moderation: A Functional Framework*, 36 BERKELEY TECH. L.J. 1091, 1095 (2021).

34. Margot E. Kaminski, *Binary Governance: Lessons from the GDPR’s Approach to Algorithmic Accountability*, 92 S. CAL. L. REV. 1529, 1533 (2019).

35. Mulligan & Bamberger, *Governance-By-Design*, *supra* note 2, at 702 (noting that “we lack a comprehensive approach—a doctrine, a set of metrics, as well as tools—for resolving design wars while accounting for the range of human rights and other public values”).

deliberatively-legitimate engagement rather than through “design wars.”³⁶ We should moreover, be cautious about “‘baking’ human and public rights values into technology systems” at any one given moment, “because of the strength and durability of a decision to govern by design.”³⁷ And when a technical component—for example a model or data set—may be used in a wide range of technical systems, and in a wide range of contexts, sound governance requires attention to the diversity of values configurations that may be required to meet regulatory goals and normative expectations. Thus, policymakers should strive for policies that “steer the protection of rights and values to the least intrusive point” to “enable the promotion of values rather than fixing them in determinatively,” and provide “technological hooks that permit different value choices in different contexts.”³⁸

B. INSTITUTIONAL DESIGN IN AI REGULATION: WHO SHOULD GOVERN?

Concerns about the fora and processes through which choices are made about what rights are prioritized in AI governance directly implicate questions of institutional design. A key component of institutional design is ensuring the appropriate expertise necessary for governance.

Accordingly, a second debate in AI governance concerns whether AI should be regulated through new, rather than existing, institutions. This debate came to the fore during a series of hearings in mid-2023 convened by the Senate Judiciary Committee, together with the Senate Homeland Security and Governmental Affairs Committee, during which senators spoke with AI researchers and industry leaders to discuss whether new developments in AI technology warrant additional regulation.³⁹ Many AI industry leaders, including Sam Altman, the CEO of OpenAI, Jared Kaplan and Jack Clark, the co-founders of Anthropic AI, and Elon Musk, the CEO of Tesla and X, also participated in a series of meetings convened by the Bipartisan Senate AI

36. *Id.*; see also Deirdre K. Mulligan & Kenneth A. Bamberger, *Apple v. FBI: Just One Battle in the ‘Design Wars’*, LAW.COM (Mar. 21, 2016), <https://www.law.com/sites/lawcomcontrib/2016/03/18/apple-v-fbi-just-one-battle-in-the-design-wars/>.

37. Mulligan & Bamberger, *Governance-By-Design*, *supra* note 2, at 750.

38. *Id.*

39. See, e.g., Michael D. Bopp, Roscoe Jones Jr., Alexander Southwell, Amanda H. Neely, Daniel P. Smith, Frances Waldmann, Kirsten Bleiweiss & Madelyn Mae La France, “*Oversight of AI: Rules for Artificial Intelligence*” and “*Artificial Intelligence in Government*” Hearings, GIBSON DUNN (June 6, 2023), <https://www.gibsondunn.com/oversight-of-ai-rules-for-artificial-intelligence-and-artificial-intelligence-in-government-hearings/> (describing takeaways from Senate Judiciary Hearings on AI).

Working Group, led by Senate Majority Leader Chuck Schumer, Senator Mike Rounds, Senator Martin Heinrich, and Senator Todd Young in 2023.⁴⁰

During these hearings and discussions, industry leaders supported the creation of new institutions to regulate AI.⁴¹ At Senate hearings convened on May 16, 2023, for example, Sam Altman of OpenAI endorsed the formation of a “new agency [for AI] that licenses any effort above a certain scale of capabilities and can take that license away and ensure compliance with safety standards.”⁴² When pressed by Senator Lindsey Graham on whether the most effective way to combat security threats stemming from AI would be to “have an agency that is more nimble and smarter than Congress” overseeing AI, Altman replied that OpenAI would be “enthusiastic” about the creation of such an agency.⁴³ When Christina Montgomery, IBM’s chief privacy and trust officer, raised doubts about whether a new agency was necessary for effective regulation of AI or whether existing institutions were sufficient, she was quickly shut down by Professor Marcus, Senator Graham, and others at the hearing who were apparently already convinced that a new federal agency to regulate AI was the right path forward.⁴⁴ Industry representatives such as Eric Schmidt, the former CEO of Google, and Mustafa Suleyman, the CEO of Microsoft AI, have also supported the creation of nonregulatory expert-led bodies to inform governments about developments in the AI space, comparable to the Intergovernmental Panel on Climate Change (IPCC).⁴⁵

40. MAJORITY LEADER CHUCK SCHUMER, SEN. MIKE ROUNDS, SEN. MARTIN HEINRICH & SEN. TODD YOUNG, BIPARTISAN SENATE AI WORKING GROUP, DRIVING U.S. INNOVATION IN ARTIFICIAL INTELLIGENCE: A ROADMAP FOR ARTIFICIAL INTELLIGENCE POLICY IN THE UNITED STATES SENATE (May 2024), https://www.schumer.senate.gov/imo/media/doc/Roadmap_Electronic1.32pm.pdf.

41. See Alexander C. Kurtz, *Regulating into the Void: Existential Uncertainty from A.I. Necessitates a New Federal Research Agency*, B.C. INTELL. PROP. & TECH. F. 1 (2024).

42. *Oversight of A.I.: Rules for Artificial Intelligence: Hearing Before the Subcomm. on Priv., Tech., and the L.*, 118th Cong. (2023), <https://www.govinfo.gov/content/pkg/CHRG-118shrg52706/html/CHRG-118shrg52706.htm>.

43. *Id.*

44. *Id.*

45. Mustafa Suleyman & Eric Schmidt, *Mustafa Suleyman and Eric Schmidt: We Need an AI Equivalent of the IPCC*, FIN. TIMES (Oct. 18, 2023), <https://www.ft.com/content/d84e91d0-ac74-4946-a21f-5f82eb4f1d2d>; Mustafa Suleyman, Mariano-Florentino (Tino) Cuéllar, Ian Bremmer, Jason Matheny, Philip Zelikow, Eric Schmidt & Dario Amodei, *Proposal for an International Panel on Artificial Intelligence (AI) Safety (IP AIS): Summary*, CARNEGIE ENDOWMENT FOR INT’L PEACE (Oct. 27, 2023), <https://carnegieendowment.org/posts/2023/10/proposal-for-an-international-panel-on-artificial-intelligence-ai-safety-ipais-summary?lang=en> (proposing an International Panel on Artificial Intelligence Safety inspired by the Intergovernmental Panel on Climate Change, a United Nations body that assesses the scientific, technical, and socio-economic information relevant to understanding climate change, its impacts, and potential risks).

Although the academic literature on AI institutions is still in its infancy, some scholars have echoed industry leaders' and policymakers' calls for the creation of a new federal agency to regulate AI in the United States, at least vis-à-vis "existential risks" such as mass access to bioweapons facilitated by generative AI (GAI) or other evolutions of AI that could be "misaligned" with human interests.⁴⁶ For example, during the same Senate hearing at which Sam Altman of OpenAI testified, Professor Gary Marcus, an expert on AI policy based at New York University, voiced a similar view to Altman's. Marcus claimed that a new, "Cabinet-level organization" with technical expertise should be propped up to regulate AI at the federal level.⁴⁷

These calls from industry leaders and academics for new agencies to regulate AI gained traction in Congress. In January 2024, Senators Michel Bennet, Elizabeth Warren, Lindsey Graham, and Peter Welch wrote a letter to Senator Schumer requesting the creation of a new federal agency to regulate digital markets including AI.⁴⁸ The Senators claimed that the hearings made evident the need to meet "the transformative challenge of AI with a thoughtful and effective regulatory framework," and voiced their belief that "this moment [in AI development] requires a new federal agency to protect consumers, promote competition, and defend the public interest."⁴⁹ The letter cited other instances—including the creation of the Food and Drug Administration in 1906 and the creation of the Federal Communications Commission in 1934—

46. See, e.g., Bryan Druzin, Anatole Boute & Michael Ramsden, *Confronting Catastrophic Risk: The International Obligation to Regulate Artificial Intelligence*, 46 MICH. J. INT'L L. 173, 182, 189, 198 (2025) (arguing that "the precautionary principle requires states to act, as waiting for conclusive scientific evidence before addressing its potential existential risk may prove too late.").

47. *Oversight of A.I.: Rules for Artificial Intelligence: Hearing Before the Subcomm. on Priv., Techn., and the L., supra* note 42. Industry leaders have changed their position, arguing for "light-weight" regulations at the federal level to limit the proliferation of state laws and no longer calling for new agencies. See, e.g., *Winning the AI Race: Strengthening U.S. Capabilities in Computing and Innovation, Hearing Before the S. Comm. on Com., Sci., & Transp.*, 119th Cong. (May 8, 2025), <https://www.govinfo.gov/content/pkg/CHRG-119shrg61426/pdf/CHRG-119shrg61426.pdf>, (responses of Sam Altman, CEO, OpenAI ("One federal framework that is light touch that we can understand and that lets us, you know, move with the speed that this moment calls for seems important and fine, but the sort of every state takes a different approach here, I think would be quite burdensome and significantly impair our ability to do what we need to do."), and Brad Smith, CEO, Microsoft ("[T]he United States needs to be in the game internationally to influence the rest of the world. And you cannot be in the game if you do nothing. You must do something. So you take . . . a lightweight approach . . . and then you build support around it.")).

48. Letter from Senators Michael F. Bennet, Lindsey O. Graham, Elizabeth Warren & Peter Welch to Chuck Schumer, Majority Leader (Jan. 23, 2024), <https://www.warren.senate.gov/imo/media/doc/F46611AE0DF77719F8B18AE6C197C52B.joint-letter-to-schumer-on-digital-platform-agency.pdf>.

49. *Id.* at 1.

where Congress, confronted with the “emergence of complex, risk-prone industries . . . elected to create [new] regulatory bodies.”⁵⁰ And the Senators justified their call for a new agency by claiming that “[p]arceling out oversight to various agencies will result in a fragmented regulatory landscape ripe for exploitation by companies with market caps greater than many countries’ gross domestic product.”⁵¹

Legal scholars and policymakers have long explored the benefits of creating new, specialized agencies to deal with novel problems,⁵² such as the ability of such institutions to develop expertise to bear on topics that existing agencies do not have the knowledge or ability to handle,⁵³ or to coordinate action among a wide range of stakeholders at the state and federal level. Such arguments about expertise underpinned the establishment of the FAA in 1958,⁵⁴ while those about coordination provided the primary rationale behind the creation of the EPA.⁵⁵

Yet many of the regulatory concerns surrounding AI are distinct from those that led to the founding of the FAA or the EPA. Whereas the invention of airplanes resulted in the advent of an entirely new sector, the same cannot be said of AI. AI has many uses and will be embedded in products and services across various sectors, but it is not a *thing* in and of itself. By nature, AI will have countless different applications and will implicate a range of rights and values.

The risk of regulatory capture⁵⁶ of specialist agencies by the interest group they regulate, moreover, is especially pronounced⁵⁷ in the AI context, because of the limited number of powerful players in the space, all of whom have nearly endless financial resources and similar regulatory interests. Compared to a generalist institution like the Office of Information and Regulatory Affairs or even an agency that regulates various industries like the EPA, the benefits of

50. *Id.* at 2.

51. *Id.*

52. See generally Rachel E. Barkow, *Insulating Agencies: Avoiding Capture Through Institutional Design*, 89 TEX. L. REV. 15 (2010).

53. See *id.* at 20.

54. See generally John W. Gelder, *Air Law - The Federal Aviation Act of 1958*, 57 MICH. L. REV. 1214 (1959).

55. See Jonathan H. Adler, *The Environmental Protection Agency Turns Fifty*, 70 CASE W. RES. L. REV. 871 (2020).

56. Capture being the phenomenon by which “organized interest groups successfully act to vindicate their goals through government policy at the expense of the public interest.” Michael A. Livermore & Richard L. Revesz, *Regulatory Review, Capture, and Agency Inaction*, 101 GEO. L.J. 1337, 1340 (2013).

57. See Jonathan R. Macey, *Organizational Design and Political Control of Administrative Agencies*, 8 J. L. ECON. & ORG. 93, 99 (1992) (“The interest group that is regulated by a single regulatory agency will be able to influence that agency to a far greater extent than the interest groups that must ‘share’ their agency with a variety of other interest groups.”).

capturing a specialized institution are comparatively higher for industry players. They have the undivided attention of regulators vis-à-vis a given regulation, they do not have to compete with other industries or interests for regulators' attention, and they have incentives to coordinate when their interests are aligned.⁵⁸ The Senators correctly identified the challenges posed by the deep pockets of industry, but the creation of a new entity to regulate a small number of companies with such wealth creates the perfect conditions for regulatory capture.

Finally, and perhaps most importantly in light of the analysis below, the assignment of AI regulation to a specialized regulator would create its own set of problems of legitimacy and expertise. For as we discuss, *infra*, the impact of AI on real interests and real individuals does not play out at a theoretical level, but on the ground, in concrete contexts that are already often governed by regulatory bodies with domain competence. While these organs of government might face challenges of technological expertise, they are already invested with legitimacy in administering laws and protecting and enforcing rights independently defined in legislation and regulation. And they possess critical expertise relevant to the ways that AI would affect rights in particular contexts. Importantly, the fields these agencies regulate have their own logics—epistemological and ethical—that control how they produce and act upon knowledge. Redirecting regulatory authority to a specialized agency would decenter that domain expertise in AI governance. Centering the tool or method—AI—rather than the domain creates the conditions for regulatory displacement, or a form of regulatory arbitrage by which firms might seek to undermine, escape, or preclude meaningful domain-specific regulation by appeals to generalized, and less tailored or exacting, mandates.

These concerns suggest that addressing regulatory expertise requires a fundamentally different approach—one that better prioritizes public values. This approach should treat AI as “normal technology”⁵⁹ that can be shaped towards different ends,⁶⁰ rather than fetishizing it as uniquely ungovernable. Effective AI regulation must be domain-specific, vindicate relevant rights, and ensure meaningful participation by affected stakeholders.

58. See generally Livermore & Revesz, *supra* note 56.

59. Arvind Narayanan & Sayash Kapoor, *AI as Normal Technology*, KNIGHT FIRST AMEND. INST. (Apr. 15, 2025), <https://knightcolumbia.org/content/ai-as-normal-technology>.

60. Langdon Winner, *Do Artifacts Have Politics?*, 109 DAEDALUS 121, 123 (1980) (describing that the politics of a technical system can arise through their design processes or for a narrow group of technologies through inherent properties). Narayanan and Kapoor's “normal technology” argument, in part, positions AI in the first category—it is pliable and configurable and can support different social and political arrangements.

III. OUR FRAMEWORK FOR RECENTERING PUBLIC VALUES IN TECHNOLOGY GOVERNANCE: AN ALTERNATIVE TO THE AI DEBATES

The framework for technology governance set forth in our earlier work⁶¹ reflects the concerns articulated by different sides in the prevailing debates over AI regulation. AI governance requires choices about, a focus on, and a comprehension of, the public rights implicated by the technology. It requires expertise sufficient to comprehend the risks posed to those rights in concrete contexts. Yet our regulatory principles suggest an alternate approach that accounts for both the limits of the risk/rights binary and the shortcomings of traditional legal and governance bodies in the face of technology challenges. In particular, they address challenges that include: inadequate technical expertise; the difficulty of accounting for multiple rights and values, or contests between them; and the absence of fora for, and meaningful assessments to inform, discussions regarding how to prioritize values in the context of technological development and implementation. Looking forward, these principles seek to ensure governance processes that prioritize rights and reflect fundamental democratic norms, are free from capture or caprice, and avoid technological fixes that privilege singular values and often disfavor human rights.

To do so, our framework sets forth three interrelated regulatory principles for centering the “public” in technology governance.

First, such governance must *privilege human and public rights*, meaning on the one hand that design must be oriented intentionally with rights in mind and, on the other, that it should reflect the reality that a multiplicity of rights might be impacted by AI systems, and that the rights impacted might depend on domain and context. Accordingly, design should allow, wherever possible, for flexibility in deployments and make values choices visible and configurable for deployers and users. This principle does not, however, reject a regulatory focus on risk or risk assessment. Rather, it requires that risk management must be responsive to rights in context and reflect the fact that the rights we care about and how we evaluate and mitigate risks to them vary by domain.

Second, AI governance requires *agencies with appropriate domain competence and sociotechnical expertise*. These are venues in which “government and other stakeholder groups have deep access to technical expertise”⁶² in addition to “domain-specific expertise.”⁶³ While it focuses on institutional design, this principle points away from debates over the specific structure of regulatory institutions to the question of whether regulators possess the right tools to

61. See discussion *supra* notes 2–10 and accompanying text.

62. Mulligan & Bamberger, *Governance-By-Design*, *supra* note 2, at 759.

63. Mulligan & Bamberger, *Procurement as Policy*, *supra* note 1, at 837.

avoid goal-myopia, foster the consideration of competing values, claim legitimacy in the choice between values, and understand the implications for them of technical decisions.⁶⁴ Such tools are required in terms of both authority and competence, including technical—or our preferred term, sociotechnical—expertise.

Finally, our third regulatory principle mandates *maintaining the publicness of policymaking*, specifically by employing “mechanisms that translate traditional commitments to participation and transparency to the technology context, in ways that address the intricate way in which policy is embedded in technical design and implementation choices.”⁶⁵

Stakeholder-participation and community engagement processes go part of the way to integrating public deliberation of values in technology policymaking. Yet, as we have argued,⁶⁶ they cannot alone address the challenge that technology design presents for governance principles of participation and transparency because of the difficulty in comprehending the value implications of technology systems from the outside. The challenge is two-fold. On the one hand, those systems’ opacity obscures the values embedded in technical infrastructure from the start; on the other, the impact on rights and values plays out not just at the moment of initial design, but “in a continuum—at design time, configuration time, and run time.”⁶⁷ Thus the impact on rights and interests varies throughout these phases in different contexts and applications.

Accordingly, we contend: (1) meaningful participation requires sociotechnical expertise among both regulators and stakeholders; and (2) meaningful transparency must involve “political visibility” into the existence and political nature of questions being resolved during design and use.⁶⁸ To that end, we have argued for the use of values-surfacing tools in technical design, drawing on a range of approaches that provide clarity over the properties embedded in code and other technical artifacts, as well as its performance.

A wide range of tools can keep values in view at different stages of development, configuration, and deployment. For example, formal methods provide clarity about a technology’s properties during development and use. Impact assessments and data and system documentation can serve as *boundary*

64. Mulligan & Bamberger, *Governance-By-Design*, *supra* note 2, at 759–69.

65. *Id.* at 770.

66. *Id.* at 770–83.

67. *Id.* at 773 (citing David D. Clark, John Wroclawski, Karen R. Sollins & Robert Braden, *Tussle in Cyberspace: Defining Tomorrow’s Internet*, 13 IEEE/ACM TRANSACTIONS ON NETWORKING 462, 463 (2005)).

68. *Id.* at 776.

objects.⁶⁹ that facilitate designers and the public's deliberation on the “technocratic and democratic elements”⁷⁰ of systems. These tools can support learning, communication and deliberation across diverse stakeholders on a shared endeavor such as the design of a sociotechnical system. For example, the Census Bureau published data artifacts (aka “demonstration data”) produced by different system implementations of differential privacy to scaffold public understanding of those embedded policy choices.⁷¹ This demonstration data served as a boundary object “allow[ing] stakeholders to interactively and intuitively explore the impact of potential implementation choices on their equities.”⁷²

As we discuss, *infra*,⁷³ identifying the methods and the metrics for meaningfully aligning a sociotechnical system with relevant values and assessing its impact on public rights and values presents challenges. Yet such tools and methods offer the capacity to both catalyze deliberation about the technical aspects of system design and also to surface the political implications of those choices. These tools and methods thus “bridge the dual deliberation requirements of substantive expertise and political visibility”⁷⁴ by creating “different frameworks and bring new considerations to bear in agency actions.”⁷⁵ And at the same time, they “bridge the gulf between the substantive domain expertise of agency staff and the frameworks and knowledge of outside experts,”⁷⁶ facilitating “participation by issue experts and by stakeholders who might otherwise be unaware of relevant risks and technological alternatives.”⁷⁷

69. Susan Leigh Star & James R. Griesemer, *Institutional Ecology, 'Translations' and Boundary Objects: Amateurs and Professionals in Berkeley's Museum of Vertebrate Zoology, 1907–39*, 19 SOC. ST. SCI. 387, 388 (1989).

70. Mulligan & Bamberger, *Procurement as Policy*, *supra* note 1, at 842.

71. Amina A. Abdu, Lauren M. Chambers, Deirdre K. Mulligan & Abigail Z. Jacobs, *Algorithmic Transparency and Participation Through the Handoff Lens: Lessons Learned from the U.S. Census Bureau's Adoption of Differential Privacy*, FACCT'24: PROCS. OF THE 2024 ACM CONF. ON FAIRNESS, ACCOUNTABILITY, & TRANSPARENCY 1150–62 (2024).

72. *Id.* at 1158.

73. *See infra* Part IV.

74. Mulligan & Bamberger, *Procurement as Policy*, *supra* note 1, at 842.

75. *Id.* at 844.

76. *Id.*

77. Mulligan & Bamberger, *Governance-By-Design*, *supra* note 2, at 765.

IV. REFLECTING OUR GOVERNANCE PRINCIPLES: EXAMPLES FROM THE BIDEN-HARRIS ADMINISTRATION'S APPROACH TO AI GOVERNANCE

The Biden-Harris Administration moved forward with a suite of AI Governance initiatives that reflect these governance principles.⁷⁸ We focus on three examples.

First, the Administration made clear that AI and automated decision-making systems—regardless of who built or used them—should privilege the public's rights—particularly human rights—and safety.⁷⁹ As the President's Science and Technology Advisor and Director of OSTP, Arati Prabhakar said: “[w]e’re in choppy waters with this rapidly changing technology, and that

78. See discussion *infra* Part IV. While Mulligan served as Principal Deputy U.S. Chief Technical Officer in the White House Office of Science and Technology Policy (OSTP) in the Biden-Harris Administration during a key period of AI policy development, this is a retrospective analysis of some, but not all, of the Administration's key actions, not a claim that our research framing specifically drove its approach. In addition, this review does not include some key Administration actions designed to limit adversarial uses of AI by foreign actors including export controls on the most powerful models, the “Diffusion Rule,” Framework for Artificial Intelligence Diffusion, 90 Fed. Reg. 4544 (Jan. 15, 2025), which revised the Export Administration Regulations’ (EAR) controls on advanced computing integrated circuits (ICs) and added controls on AI model weights for certain advanced closed-weight dual-use AI models to protect U.S. national security and foreign policy interests among other things, rescinded three days before its effective date on May 12, 2025. See *Department of Commerce Rescinds Biden-Era Artificial Intelligence Diffusion Rule, Strengthens Chip-Related Export Controls*, BUR. OF INDUS. & SEC., U.S. DEP’T OF COM. (May 12, 2025), <https://media.bis.gov/sites/default/files/documents/05.07%20Recission%20of%20AI%20Diffusion%20Press%20Release-2.pdf>; and the “BIS Rule,” issued by the Department of Commerce Bureau of Industry and Standards pursuant to the Defense Production Act of 1950 (DPA), which requires quarterly reporting from companies developing or demonstrating an intent to develop dual-use foundation AI models and those with large-scale computing clusters for AI model training runs over 10²⁶ computation operations or acquiring/possessing a computer cluster with data center networking over 300 Gbits. Establishment of Reporting Requirements for the Development of Advanced Artificial Intelligence Models and Computing Clusters, 89 Fed. Reg. 73612 (proposed Sep. 11, 2024) (to be codified at 15 C.F.R. pt. 702). This rule has not yet been finalized, and it seems unlikely that it will be given recent Administration actions. During the Biden-Harris Administration, BIS required several AI companies to make such disclosures pursuant to E.O. 14110.

79. Exec. Order No. 14,110 § 2, 88 Fed. Reg. 75191, at 75191–93 (Nov. 1, 2023) (signed Oct. 30, 2023) (setting out the principles and policies to guide AI including: § 2(a) ensuring that AI is safe and secure; § 2(b) promoting responsible innovation, competition and collaboration; § 2(c) supporting workers; § 2(d) advancing equity and civil rights; § 2(e) protecting consumers; § 2(f) protecting privacy and civil liberties; § 2(g) building the necessary expertise—technical, managerial, ethical, legal, etc.—in government; § 2(h) leading global development and adoption so that AI “benefits the whole world, rather than exacerbating inequities, threatening human rights, and causing other harms”). This Executive Order was revoked by the Trump administration on January 20, 2025.

means it's more important than ever to steer by the light of these fundamental values.”⁸⁰ At the first AI summit, Vice President Harris stated that:

[T]he urgency of this moment must then compel us to create a collective vision of what this future must be. A future where AI is used to advance human rights and human dignity, where privacy is protected and people have equal access to opportunity, where we make our democracies stronger and our world safer.⁸¹

Harris further stated that the Administration “believe[d] that all leaders from government, civil society, and the private sector have a moral, ethical, and societal duty to make sure that AI is adopted and advanced in a way that protects the public from potential harm and that ensures that everyone is able to enjoy its benefits.”⁸² Even in the context of national security, the President centered rights and values, writing that:

Success for the United States in the age of AI will be measured not only by the preeminence of United States technology and innovation, but also by the United States’ leadership in developing effective global norms and engaging in institutions rooted in international law, human rights, civil rights, and democratic values.⁸³

To protect the public’s rights and safety the Administration took a layered approach to governance, empowering domain specific regulators and creating a new set of expertise to develop risk identification and mitigation techniques.

Second, the Administration expanded the AI and AI-enabling expertise within the government; clarified the importance of civil rights, privacy civil liberties, and other rights-centered experts in AI design and use; and created opportunities for AI experts in government, academia, civil society, and the corporate sector to collaboratively build a body of knowledge and practice to support the identification and development of testing, evaluation, validation

80. OSTP Director Arati Prabhakar Remarks on Managing AI’s Risks to Seize its Benefits, as Prepared for Delivery at the *Carnegie Endowment for International Peace*, WHITE HOUSE (Nov. 14, 2023), <https://bidenwhitehouse.archives.gov/ostp/news-updates/2023/11/14/remarks-of-arati-prabhakar-at-carnegie-endowment-for-international-peace/>.

81. Vice President Kamala Harris, Remarks on the Future of Artificial Intelligence at the U.S. Embassy, London, United Kingdom (Nov. 1, 2023), <https://bidenwhitehouse.archives.gov/briefing-room/speeches-remarks/2023/11/01/remarks-by-vice-president-harris-on-the-future-of-artificial-intelligence-london-united-kingdom/>.

82. *Id.*

83. Memorandum on Advancing the United States’ Leadership in Artificial Intelligence; Harnessing Artificial Intelligence to Fulfill National Security Objectives; and Fostering the Safety, Security, and Trustworthiness of Artificial Intelligence, 2024 DAILY COMP. PRES. DOC. 00945, 3 (Oct. 24, 2024).

and risk management practices to support federal agencies and others developing and using, as well as regulating, AI systems.⁸⁴

Third, the Administration instituted practices to expose and interrogate the values embedded in federal agency AI *use cases* throughout design and deployment, and scaffold public participation in their design, risk mitigation activities, and evaluations.⁸⁵

Together, these initiatives reflect the interdependent principles of our governance-by-design framework, chart a new path in AI governance, and diverge from many of the dominant AI debates. The AI governance framework is not a traditional risk management framework but rather layers new knowledge onto and offers new expertise to support sectoral regulations and regulators that define rights and determine the level of tolerable risk within domains. It offers a model by which AI governance can be both rights and risk based, and can leverage new and existing institutions effectively, all while engaging key stakeholders from the public and private sectors, civil society, and academia, at every stage from AI testing to deployment.

A. PRIVILEGING HUMAN AND PUBLIC RIGHTS: MAINTAINING EXPERT AGENCY AUTHORITY WHILE BUILDING A SHARED KNOWLEDGE BASE FOR RISK ASSESSMENT METHODS AND PRACTICES

As a foundational matter, the Biden-Harris Administration made it clear from the outset that, regardless of the parties involved in their development and utilization, AI and automated decision-making systems must be developed and implemented in ways that account for generalized and sector-specific systemic risks, but that the fundamental priority should be mitigating risks to public rights, with a particular emphasis on human rights and safety.⁸⁶

The Administration issued the Blueprint for an AI Bill of Rights (AI BoR), a clear statement affirming the centrality of the public's rights and safety in the design and use of AI and articulating processes and practices to protect them.⁸⁷ The White House OSTP published the AI BoR in October 2022, before

84. Discussed, *infra*, at Section IV(B).

85. Discussed, *infra*, at Section IV(C).

86. Below we focus on a subset of the Administration's AI actions to illustrate the focus on rights and safety. Other actions to limit risk include the AI diffusion rule, subsequently withdrawn under the Trump Administration, and the BIS reporting rule. Framework for Artificial Intelligence Diffusion, 90 Fed. Reg. 4544 (Jan. 15, 2025); *Department of Commerce Announces Rescission of Biden-Era Artificial Intelligence Diffusion Rule, Strengthens Chip-Related Export Controls*, BUR. OF INDUS. & SEC., U.S. DEP'T OF COM. (May 13, 2025), <https://www.bis.gov/press-release/departments-commerce-announces-rescission-biden-era-artificial-intelligence-diffusion-rule-strengthens>.

87. WHITE HOUSE OFF. OF SCI. & TECH. POL'Y, BLUEPRINT FOR AN AI BILL OF RIGHTS: MAKING AUTOMATED SYSTEMS WORK FOR THE AMERICAN PEOPLE (2022), <https://bidenwhitehouse.archives.gov/ostp/ai-bill-of-rights/> [hereinafter AI BOR].

ChatGPT and other large language models garnered the attention of the media, the public, and policymakers.⁸⁸ That document set out a framework to guide the design, development, and deployment of automated systems so that they protect the rights of the American public and reinforce our nation's longstanding values. The AI BoR announced five principles to guide the design, use, and deployment of automated systems.⁸⁹ Automated decision-making systems (ADS), including AI systems, should be safe and effective, free from algorithmic discrimination, and respect data privacy.⁹⁰ Individuals should know when ADS are being used and receive an explanation of how it impacts decisions about them.⁹¹ Lastly, individuals should be able to opt for a human alternative rather than an ADS process and have easy access to a person who can quickly consider and address problems with ADS systems.⁹² The AI BoR pays special attention to domains, including criminal justice and education, where automated systems can have significant adverse effects on human rights, civil liberties, and civil rights.⁹³ It calls for limitations on surveillance, including stating that continuous surveillance and monitoring should not be used in settings where it is likely to limit rights, opportunities, or access.⁹⁴

The Administration quickly moved from rights-based principles to commitments to protection. First, in February 2023, President Biden directed the federal government to root out bias in the design and use of new technologies, such as artificial intelligence; to protect the public from algorithmic discrimination; and to bring civil rights offices into conversations about the design, procurement, and use of automated systems.⁹⁵

Second, the Administration pushed the private sector to protect the public's rights and safety during AI development and use. In July 2023, President Biden garnered voluntary commitments from leading AI companies, including Amazon, Anthropic, Google, Inflection, Meta, Microsoft, and OpenAI, to engage in a set of practices designed to identify and mitigate risks to the public's rights and safety.⁹⁶ Specifically, these companies committed to

88. See *When Was ChatGPT Released*, SCRIBBR, <https://www.scribbr.com/frequently-asked-questions/when-was-chatgpt-released/> (last visited Sep. 16, 2025) (noting that Chat GPT was publicly released on November 30, 2022, and that at the time of its release, was described as a "research preview").

89. AI BOR, *supra* note 87, at 5–7.

90. *Id.* at 5–6.

91. *Id.* at 6.

92. *Id.* at 7.

93. See, e.g., *id.* at 36.

94. *Id.* at 6, 30, 34.

95. Exec. Order No. 14,091, 88 Fed. Reg. 10825 (Feb. 16, 2023).

96. *FACT SHEET: Biden-Harris Administration Secures Voluntary Commitments from Leading Artificial Intelligence Companies to Manage the Risks Posed by AI*, WHITE HOUSE (July 21, 2023), <https://bidenwhitehouse.archives.gov/briefing-room/statements-releases/2023/07/21/>

“internal and external security testing of their AI systems before their release”; “sharing information across the industry and with governments, civil society, and academia on managing AI risks”; and “prioritizing research on the societal risks that AI systems can pose, including on avoiding harmful bias and discrimination, and protecting privacy,” among other commitments.⁹⁷

Third, in October 2023, President Biden issued an Executive Order on Safe, Secure, and Trustworthy Artificial Intelligence (EO 14110).⁹⁸ This executive order directed sweeping actions across the government. It opens by affirming the values—privacy, equity, etc.—that AI design and use must respect and mitigate its impact on (including by specifically stating that “Americans’ privacy and civil liberties must be protected as AI continues advancing”).⁹⁹ EO 14110 specifically notes that while the Administration aimed to “promote responsible uses of AI that protect consumers, raise the quality of goods and services, lower their prices, or expand selection and availability,” it would also ensure protections for important rights “especially important in critical fields like healthcare, financial services, education, housing, law, and transportation, where mistakes by or misuse of AI could harm patients, cost consumers or small businesses, or jeopardize safety or rights.”¹⁰⁰ Through EO 14110, the Administration encouraged existing agencies to make full use of their authorities to address domain-relevant AI risks and, as we discuss, *infra*, set up a hiring surge and a new institution to expand the AI and AI-enabling expertise available across federal agencies.

In March and July 2024, respectively, the Office of Management and Budget (OMB) issued guidance to federal agencies on the responsible use of AI.¹⁰¹ In October 2024, the White House issued a national security

fact-sheet-biden-harris-administration-secures-voluntary-commitments-from-leading-artificial-intelligence-companies-to-manage-the-risks-posed-by-ai/.

97. *FACT SHEET: Biden-Harris Administration Secures Voluntary Commitments from Eight Additional Artificial Intelligence Companies to Manage the Risks Posed by AI*, WHITE HOUSE (Sep. 12, 2023), <https://bidenwhitehouse.archives.gov/briefing-room/statements-releases/2023/09/12/fact-sheet-biden-harris-administration-secures-voluntary-commitments-from-eight-additional-artificial-intelligence-companies-to-manage-the-risks-posed-by-ai/#:~:text=The%20companies%20commit%20to%20internal,circumvent%20safeguards%2C%20and%20technical%20collaboration.>

98. Exec. Order No. 14,110, 88 Fed. Reg. 75191 (Nov. 1, 2023) (signed Oct. 30, 2023).

99. Exec. Order No. 14,110 § 1(f), 88 Fed. Reg. at 75193.

100. Exec. Order No. 14,110 § 1(e), 88 Fed. Reg. at 75193.

101. OFF. OF MGMT. & BUDGET, EXEC. OFF. OF THE PRESIDENT, MEMO. NO. M-24-10, ADVANCING GOVERNANCE, INNOVATION, AND RISK MANAGEMENT FOR AGENCY USE OF ARTIFICIAL INTELLIGENCE (Mar. 28, 2024), <https://www.whitehouse.gov/wp-content/uploads/2024/03/M-24-10-Advancing-Governance-Innovation-and-Risk-Management-for-Agency-Use-of-Artificial-Intelligence.pdf> [hereinafter OMB MEMO. M-24-10]. Executive Order 14110 directed OMB to fulfill unmet obligations under 40 U.S.C. § 11301 (the AI in Government Act and the Advancing American AI Act).

memorandum (NSM) and accompanying framework establishing similar requirements for agencies on use of AI on national security systems.¹⁰² The OMB guidance directs agencies to ensure the public's rights and interests—including privacy, nondiscrimination and equity, security, and accessibility—are protected and risks to them are mitigated in system design and use.¹⁰³ OMB's guidance to agencies established the most specific and rigorous set of requirements for the design and use of technology by government agencies. The NSM is similarly robust and includes prohibitions on certain uses of AI.¹⁰⁴ OMB subsequently developed AI specific procurement policy to ensure agencies received information and secured the ability to interact with AI systems necessary to evaluate the impact of federal AI use cases on the public's rights and safety.¹⁰⁵

The Biden-Harris Administration affirmed the centrality of human rights and democratic values in multinational AI policy as well. In November 2023, Vice President Kamala Harris participated in a Global Summit on AI Safety hosted by former Prime Minister Rishi Sunak of the United Kingdom.¹⁰⁶ Prior to this engagement, Vice President Harris laid out the Administration's vision of a future where "AI is used to advance human rights and human dignity, where privacy is protected and people have equal access to opportunity,"¹⁰⁷ and announced the Administration's joint commitment with thirty other countries on the responsible use of military AI.¹⁰⁸ In her speech, she pushed back against an exclusive emphasis on AI risk related to chemical, biological, radiological and nuclear risks, stating:

But let us be clear. There are additional threats that also demand our action—threats that are currently causing harm and which, to many people, also feel existential. Consider, for example: When a senior is kicked off his healthcare plan because of a faulty AI algorithm, is

102. WHITE HOUSE, FRAMEWORK TO ADVANCE AI GOVERNANCE AND RISK MANAGEMENT IN NATIONAL SECURITY 1–2 (Oct. 24, 2024).

103. OMB MEMO. M-24-10, *supra* note 101.

104. 2024 DAILY COMP. PRES. DOC. 00945, *supra* note 83; WHITE HOUSE, *supra* note 102.

105. OFF. OF MGMT. & BUDGET, EXEC. OFF. OF THE PRESIDENT, MEMO. NO. M-24-18, ADVANCING THE RESPONSIBLE ACQUISITION OF ARTIFICIAL INTELLIGENCE IN GOVERNMENT (Sep. 24, 2024), <https://www.whitehouse.gov/wp-content/uploads/2024/10/M-24-18-AI-Acquisition-Memorandum.pdf> [hereinafter OMB MEMO. M-24-18].

106. Press Release, Statement by Press Secretary Kirsten Allen on the Vice President's and Second Gentleman's Travel to the United Kingdom (Oct. 26, 2023) (stating that the Vice President would deliver a speech on AI on November 1 in London and represent the United States at the Global Summit on AI Safety at Bletchley Park on November 2), <https://bidenwhitehouse.archives.gov/briefing-room/statements-releases/2023/10/26/statement-by-press-secretary-kirsten-allen-on-the-vice-presidents-and-second-gentlemans-travel-to-the-united-kingdom/>.

107. Vice President Kamala Harris, *supra* note 81.

108. *Id.*

that not existential for him? When a woman is threatened by an abusive partner with explicit, deep-fake photographs, is that not existential for her? When a young father is wrongfully imprisoned because of biased AI facial recognition, is that not existential for his family? And when people around the world cannot discern fact from fiction because of a flood of AI-enabled mis- and disinformation, I ask, is that not existential for democracy? Accordingly, to define AI safety, I offer that we must consider and address the full spectrum of AI risk—threats to humanity as a whole, as well as threats to individuals, communities, to our institutions, and to our most vulnerable populations.¹⁰⁹

The Administration ensured that international conversations—from those in the UN General Assembly to those held by the G7—centered human rights. The United States, for example, led the drafting and adoption of the *United Nations General Assembly resolution on trustworthy AI for sustainable development, and other international efforts*. President Biden, moreover, used his last address before the UN General Assembly to urge world leaders to “ensure that AI supports, rather than undermines, the core principles that human life has value and all humans deserve dignity.”¹¹⁰

At the same time that the Biden-Harris Administration was articulating the public rights and values at the center of its AI governance approach, it was also building the science base and practices necessary for agencies and others to test and evaluate risks associated with AI models and systems. Notably, in January 2023, the National Institute of Standards and Technology (NIST) published the Artificial Intelligence Risk Management Framework (AI RMF).¹¹¹ The AI RMF’s goal is to “offer a resource to the organizations designing, developing, deploying, or using AI systems to help manage the many risks of AI and promote trustworthy and responsible development and use of AI systems.”¹¹² Through the AI RMF, the Administration set out processes that could assist organizations in proactively and systemically identifying and limiting risks—to human rights, safety, as well as the climate—posed by AI systems. It is “voluntary, rights-preserving, non-sector-specific,

109. *Id.*

110. Ja’han Jones, *Biden Warns Dictators Could Use AI to Put ‘Shackles’ on the ‘Human Spirit’*, MSNBC (Sep. 25, 2024), <https://www.msnbc.com/the-reidout/reidout-blog/biden-un-speech-ai-rcna172702>.

111. NAT’L INST. OF STANDARDS & TECH., U.S. DEP’T OF COM., ARTIFICIAL INTELLIGENCE RISK MANAGEMENT FRAMEWORK (AI RMF 1.0) (Jan. 2023), <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf> [hereinafter AI RMF]. This was pursuant to the National Artificial Intelligence Initiative Act of 2020, authorized in division E of the William M. (Mac) Thornberry National Defense Authorization Act for Fiscal Year 2021, Pub. L. No. 116–283, 134 Stat. 3388, which passed with a Congressional override of President Trump’s veto.

112. AI RMF, *supra* note 111, at 2.

and use-case agnostic,”¹¹³ and therefore useful to a broad swath of companies, organizations, and federal agencies operating across various sectors of society. It does not disturb existing legal obligations—whether under civil rights, consumer protection, or environmental law—but rather provides process guidance to aid entities attempting to proactively build and deploy systems that reduce the possibility of interfering with rights or obligations, or other objectives an entity might independently seek to achieve.¹¹⁴ The AI RMF “core” described below also provides useful insights for regulatory and enforcement agencies of the practices for building and deploying AI, identified through an inclusive, multistakeholder process. Such non-binding process standards can inform agencies’ regulations, enforcement actions, and remedies.

The AI RMF is divided into two parts: Part I sets the table, describing NIST’s view of the risks associated with AI, defining the intended audience for the Framework, and outlining the “characteristics of trustworthy AI systems”;¹¹⁵ and Part II, the “core” of the AI RMF, describes four specific functions (namely, “GOVERN, MAP, MEASURE, AND MANAGE”) that the AI RMF’s audience can use to address AI risks and improve AI safety and trustworthiness.¹¹⁶ Notably, Part I defines risk as “the composite measure of an event’s probability of occurring and the magnitude or degree of the consequences of the corresponding event,”¹¹⁷ risk management as “coordinated activities to direct and control an organization with regard to risk,”¹¹⁸ and posits that the AI RMF offers risk management approaches that both “minimize anticipated negative impacts of AI systems *and* identify opportunities to maximize positive impacts.”¹¹⁹ It further identifies specific

113. *Id.*

114. While some have questioned the use of risk regulation techniques to address AI’s impact on rights, that critique is premised on the belief that risks regulation is displacing rights regulation. At least with respect to the AI RMF and recent OMB guidance, this is not the case. The AI RMF is voluntary and non-binding on the private sector, and while OMB MEMO. M-24-10, *supra* note 101, at 16, encouraged agencies to incorporate “additional best practices . . . from the National Institute of Standards and Technology (NIST) AI Risk Management Framework” as well as other sources including the AI BOR, and the revamped OMB guidance on the subject, OMB M-Memo 25-21 incorporates key aspects of it, it does not—and cannot—disturb existing law. The encouragement to adopt practices to systemically and proactively address risks to rights is a complement to U.S. protections for rights not a substitute framework. See Margot E. Kaminski, *Regulating the Risks of AI*, 103 B.U. L. REV. 1347, 1378 (2023) (questioning the utility of risk regulation to “address big, often-unquantifiable, often-contested, often-contextual, and often individualized ‘risks’”).

115. AI RMF, *supra* note 111, at 2.

116. *Id.* at 3.

117. *Id.* at 4.

118. *Id.*

119. *Id.*

harms to people, organizations, and ecosystems that AI risk management can prevent.¹²⁰ Importantly, the AI RMF “does not prescribe risk tolerance” nor which rights or values to prioritize—rather, it places the onus on implementing organizations to make these determinations, based on contextual and case-or sector-specific factors—informed by regulations, the needs and ethical obligations of relevant professionals, and the needs of affected communities—that can change over time.¹²¹

Pursuant to EO 14110, described, *supra*, NIST took several additional actions to assist agencies and other organizations in mitigating the risks associated with AI.¹²² President Biden established the AI Safety Institute within NIST, which was charged with advancing the science of AI safety and operationalizing capabilities testing on foundational AI models.¹²³ Other NIST actions included issuing “Guidelines for Evaluating Differential Privacy Guarantees.”¹²⁴ NIST defines differential privacy as “a privacy-enhancing technology that quantifies privacy risk to individuals when their information appears in a dataset.”¹²⁵ The guidelines were intended to “help agencies and practitioners of all backgrounds . . . better understand how to evaluate promises made (and not made) when deploying differential privacy, including for privacy-preserving machine learning.”¹²⁶ The guidelines were published alongside a supplemental interactive software archive that “illustrate[s] how to achieve differential privacy and other concepts described in the publication.”¹²⁷

120. *Id.* at 5.

121. *Id.* at 7.

122. Exec. Order No. 14,110 §§ 4.1, 4.4.b.ii, 9.b., 10.1.d.i, 88 Fed. Reg. 75191, 75196, 75201, 75217, 75219 (Nov. 1, 2023) (signed Oct. 30, 2023).

123. *FACT SHEET: Vice President Harris Announces New U.S. Initiatives to Advance the Safe and Responsible Use of Artificial Intelligence*, WHITE HOUSE (Nov. 1, 2023), <https://bidenwhitehouse.archives.gov/briefing-room/statements-releases/2023/11/01/fact-sheet-vice-president-harris-announces-new-u-s-initiatives-to-advance-the-safe-and-responsible-use-of-artificial-intelligence/> (reporting that the Vice President is announcing the creation of The United States AI Safety Institute (US AISI) inside NIST). For an overview of AISI, see U.S. A.I. SAFETY INST., NAT’L INST. OF STANDARDS & TECH., *THE UNITED STATES ARTIFICIAL INTELLIGENCE SAFETY INSTITUTE: VISION, MISSION, AND STRATEGIC GOALS* (May 21, 2024), <https://www.nist.gov/system/files/documents/2024/05/21/AISI-vision-21May2024.pdf>.

124. Joseph Near & David Darais, *NIST SP 800-226 (Initial Public Draft) Guidelines for Evaluating Differential Privacy Guarantees*, NAT’L INST. OF STANDARDS & TECH. (Dec. 11, 2023), <https://csrc.nist.gov/pubs/sp/800/226/ipd>.

125. *Id.* at 1.

126. *Id.*

127. *Id.* (describing Python Jupiter Notebook packages that illustrate how to achieve differential privacy and other concepts described in the publication); *see also* NAT’L INST. OF STANDARDS & TECH., U.S. DEP’T OF COM., *NIST-SP-800-226-SupplementalMaterial*, GITHUB (Dec. 11, 2023), <https://github.com/usnistgov/PrivacyEngCollabSpace/tree/master/tools/de-identification/NIST-SP-800-226-SupplementalMaterial/>.

NIST also published the “Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile,” a companion resource to the AI RMF.¹²⁸ This profile is an implementation of the Framework’s functions, categories, and subcategories in the context of GAI, and intends to assist organizations in “deciding how to best manage AI risks in a manner that is well-aligned with their goals, consider[ing] legal/regulatory requirements and best practices, and reflect[ing] risk management priorities.”¹²⁹ It describes the time scale, scope, and potential sources of risks from GAI, but—like the Framework itself—leaves it to organizations to determine how best to measure and tolerate GAI risks.¹³⁰ And in November 2024, NIST released a report on “existing standards, tools, methods, and practices”¹³¹ for:

authenticating [synthetic] content and tracking its provenance; labeling synthetic content, such as using watermarking; detecting synthetic content; preventing generative AI (GAI) from producing child sexual abuse material or producing non-consensual intimate imagery of real individuals (to include intimate digital depictions of the body or body parts of an identifiable individual); testing software used for the above purposes; and auditing and maintaining synthetic content.¹³²

The report outlines potential risks and harms associated with synthetic content, including “the target audience for the content; the context in which content is used or misused; the sophistication of the actor creating and/or disseminating the content; and any social, economic, and health-related (including mental health) costs” associated with its dissemination.¹³³

The Biden-Harris Administration complemented this non-exhaustive list of guidance and tooling targeted toward GAI risk management with regulations aimed at both protecting rights and mitigating risks in particular sectors. One particularly illustrative example is the Department of Health and Human Services’ (HHS) July 2024 rule on bias in AI systems.¹³⁴ Promulgated

128. NAT’L INST. OF STANDARDS & TECH., U.S. DEP’T OF COM., ARTIFICIAL INTELLIGENCE RISK MANAGEMENT FRAMEWORK: GENERATIVE ARTIFICIAL INTELLIGENCE PROFILE 1 (July 2024), <https://doi.org/10.6028/NIST.AI.600-1>.

129. *Id.* at 1.

130. *Id.* (explaining that the profile is designed to assist organizations in managing AI risks in a manner that is aligned with the organization’s risk tolerance, and resources).

131. NAT’L INST. STANDARDS & TECH., U.S. DEP’T OF COM., REDUCING RISKS POSED BY SYNTHETIC CONTENT AN OVERVIEW OF TECHNICAL APPROACHES TO DIGITAL CONTENT TRANSPARENCY 1 (Nov. 20, 2024), <https://www.nist.gov/publications/reducing-risks-posed-synthetic-content-overview-technical-approaches-digital-content>.

132. *Id.*

133. *Id.* at 2.

134. Nondiscrimination in Health Programs and Activities, 89 Fed. Reg. 37522 (May 6, 2024) (codified as amended at 42 U.S.C. § 18116).

pursuant to Section 1557 of the Affordable Care Act, this regulation aimed to combat bias in the AI systems used by healthcare providers and insurers for clinical care and health administration.¹³⁵ Section 1557 of the Affordable Care Act prohibits bias on the basis of “race, color, national origin, sex, age, or disability,”¹³⁶ and the new rule requires that covered entities take “reasonable steps” to “identify and mitigate potential discrimination” from any “augmented decision-making tools or models, such as artificial intelligence (AI) and machine learning” deployed in a health care or health insurance setting.¹³⁷ If covered entities do not take such “reasonable steps,” HHS may take “corrective action” against them.¹³⁸

This HHS rule exemplifies the Biden-Harris Administration’s layered, complementary approach to AI regulation. The Administration, in the many executive orders, memos, and public statements referenced, *supra*, made clear that equity, human rights, and democratic values were at the core of its approach to AI governance. The Administration highlighted sectors like healthcare and education where laws offered protection, but regulations needed updating to ensure that the public’s rights and safety would be protected. It laid out frameworks and practices for mitigating risks associated with AI—including privacy, generative AI, and synthetic content risks that would affect multiple sectors—but maintained the regulatory and enforcement authority of existing expert agencies. The frameworks and tools help regulated and unregulated entities mitigate risks, but do not prescriptively dictate their activities. It directed expert federal agencies to issue specific regulations that spell the ways to address AI risks to important rights at the sector level, such as the right to nondiscrimination in healthcare, and required regulated entities to mitigate risks to those rights.

In essence, the Biden-Harris Administration’s regulatory approach to AI was neither “risk-based” or “rights-based,” but incorporated principles from both approaches. It fostered the development of specialized frameworks and methods for evaluating AI risks, but because the rights at stake vary across sectors, the Administration empowered sector-specific agencies to establish how those approaches should be used to address AI-specific risks, such as the right to nondiscrimination in healthcare. Maintaining the emphasis on the substantive domains of use centers rights-expertise and positions technical experts and knowledge as a facilitator of the government’s field-specific agencies responsible for protecting rights and safety. This “field-centric” approach to AI governance maintains the centrality of the agencies’ missions,

135. *Id.* at 37522.

136. *Id.*

137. *Id.* at 37524, 37642.

138. *Id.* at 37524, 37557.

rather than emphasizing new technology entering the regulated market. It avoids technosolutionism—constructing complex phenomena as *problems* that technology is best able to solve.¹³⁹ And it corrects the misguided assumption that the use of AI reduces the relevance of existing regulations. Just as companies' use of data analytics, data mining, and statistical models in regulated areas must comply with the law, so too must the use of AI.

B. BRINGING EXPERTISE AND CAPACITY INTO GOVERNMENT: DIRECT
HIRING AND STAKEHOLDER INVOLVEMENT

The Administration used the primacy of existing domain specific regulatory and enforcement agencies and addressed the need for specialized expertise. As this section describes, EO 14110 took several steps to ensure federal agencies had the AI and AI-enabling talent to design and use AI, and to effectively regulate its use. This included upskilling staff across agencies, bringing new technical professionals into key service delivery and enforcement agencies, and creating the AI Safety Institute within the National Institutes of Standards and Technology with new AI specific staff. In addition, the Administration helped scaffold the participation of a broad range of stakeholders in AI governance activities and broadened financial investments in the field of public interest technology to bolster the expertise in civil society organizations, train future generations of public interest and mission-oriented technologists, and support research to protect the public's rights and safety.

1. *Bringing AI and AI Enabling Talent into Federal Service*

The AI and Tech Talent Task Force, created by the President through EO 14110, launched an AI Talent Surge to accelerate hiring AI and AI-enabling professionals across the federal government.¹⁴⁰ This effort included: flexible hiring authorities for federal agencies to bring in AI talent, including direct hire authorities and excepted service authorities;¹⁴¹ an interagency working group

139. Technosolutionism is both the mindset or belief that complex societal problems can be solved by technology and the construction of phenomena or things as problems which technology is best situated to solve. EVGENY MOROZOV, TO SAVE EVERYTHING, CLICK HERE: THE FOLLY OF TECHNOLOGICAL SOLUTIONISM 6 (2013) (“Solutionism . . . is not just a fancy way of saying that for someone with a hammer, everything looks like a nail; it’s not just another riff on the inapplicability of ‘technological fixes’ to ‘wicked problems’ . . . It’s not only that many problems are not suited to the quick-and-easy solutionist tool kit. It’s also that what many solutionists presume to be ‘problems’ in need of solving are not problems at all.”).

140. AI & TECH TALENT TASK FORCE, INCREASING AI CAPACITY ACROSS THE FEDERAL GOVERNMENT: AI TALENT SURGE PROGRESS AND RECOMMENDATIONS (Apr. 26, 2024), https://digital.library.unt.edu/ark:/67531/metadc2349490/m2/1/high_res_d/AI-Talent-Surge-Progress-Report.pdf.

141. Memorandum from Kiran A. Ahuja, Dir., U.S. Off. Personnel Mgmt., on Government-Wide Hiring Authorities for Advancing Federal Government Use of Artificial Intelligence (AI) to Heads of Departments and Agencies 1 (Dec. 29, 2023), <https://>

of human resources professionals, recruiting experts, technical leads, and hiring managers to share best practices on federal government-wide hiring of people with AI and other technical skills;¹⁴² guidance to agencies on how to assess and expand the AI competencies of their staff through recruitment, selection, and hiring;¹⁴³ and guidance on pay flexibility, incentive pay, and leave and workforce flexibility programs.¹⁴⁴ It also scaled up the use of government-wide tech talent programs, including the Presidential Innovation Fellows, U.S. Digital Corps, and U.S. Digital Service and the creation of the new DHS AI Corps.¹⁴⁵ Through these programs, the Administration made over 200 hires by the end of July 2024.¹⁴⁶ The White House AI and Tech Talent Task Force coordinated the processes and oversaw the distribution of many of the new professionals across federal agencies.

The Administration explicitly provided guidance that defined the skills and competencies required in the AI workforce. Importantly, that guidance clarified that key technical competencies included “[s]ociotechnical [s]ystems,” “[t]esting and [v]alidation,” and “[v]alues-[d]riven [d]esign.”¹⁴⁷ It defined “values-driven design” as:

[S]ystematically applies principles and techniques from relevant subject matter domains to all aspects of design, development, maintenance, and deployment to protect the rights and safety of stakeholders and the public, ensuring equity, security, privacy, autonomy, accessibility, justice, beneficence, and nonmaleficence. Creatively combines technical and policy approaches to protect and support these core values. [And e]nsures that values inform the design, deployment, testing, and oversight of AI systems, and that

[www.opm.gov/chcoc/transmittals/2023/Government-wide%20Hiring%20Authorities%20for%20Advancing%20Government%20Use%20of%20Artificial%20Intelligence%20\(AI\)%2012-29-2023.pdf](https://www.opm.gov/chcoc/transmittals/2023/Government-wide%20Hiring%20Authorities%20for%20Advancing%20Government%20Use%20of%20Artificial%20Intelligence%20(AI)%2012-29-2023.pdf).

142. AI & TECH TALENT FORCE, *supra* note 140, at 6.

143. Ahuja, *supra* note 141, at 1.

144. Memorandum from Kiran A. Ahuja, Dir., U.S. Off. Personnel Mgmt., on Skills-Based Hiring Guidance and Competency Model for Artificial Intelligence Work to Heads of Executive Departments and Agencies (Apr. 29, 2024), <https://www.opm.gov/chcoc/transmittals/2024/Skills-Based%20Hiring%20Guidance%20and%20Competency%20Model%20for%20Artificial%20Intelligence%20Work.pdf>.

145. AI & TECH TALENT FORCE, *supra* note 140, at 3, 6.

146. *FACT SHEET: Biden-Harris Administration Announces New AI Actions and Receives Additional Major Voluntary Commitment on AI*, WHITE HOUSE 7 (July 26, 2024), <https://bidenwhitehouse.archives.gov/briefing-room/statements-releases/2024/07/26/fact-sheet-biden-harris-administration-announces-new-ai-actions-and-receives-additional-major-voluntary-commitment-on-ai/>.

147. Ahuja, *supra* note 144, at 3.

important value-related design choices are communicated to end users.¹⁴⁸

This standard signaled an expectation that the technical workforce enlisted to build the government's capacity to use and regulate AI would have an understanding of the politics of technical artifacts¹⁴⁹ and how to construct and use them to protect core rights and values,¹⁵⁰ in addition to more standard competencies such as data visualization and machine learning. This guidance assisted agencies in identifying key skills and competencies needed for AI professionals, increased opportunities for individuals with nontraditional academic backgrounds,¹⁵¹ and broadened agencies' and the public's understanding of the skills necessary for responsible development of AI.

The Task Force worked to bolster the technical expertise at agencies attempting important AI-use cases and those with significant enforcement missions and short on technical expertise. Some agencies had funding to hire additional staff using direct hire authorities. Most notably, the Department of Homeland Security brought on nearly fifty new AI and AI-enabling personnel to work on missions, such as countering fentanyl networks, combating child sexual exploitation and abuse, and delivering immigration services.¹⁵² In addition, the Administration launched a novel effort to build a pool of Science, Technology, Engineering, and Math (STEM) and AI experts that could be flexibly deployed to help agencies support implementation of the EO 14110, as well as the National Security Memorandum on Revitalizing America's

148. *Id.* at 16.

149. Winner, *supra* note 60, at 121–36.

150. Katie Shilton, Jes A. Koepfler & Kenneth R. Fleischmann, *How to See Values in Social Computing: Methods for Studying Values Dimensions*, CSCW '14: PROCS. OF THE 17TH ACM COMPUT. SUPPORTED COOP. WORK & SOC. COMPUTING, 426–35 (2014); Cory Knobel & Geoffrey C. Bowker, *Values in Design*, 54 COMMUN. ACM 26 (2011); Mary Flanagan, Daniel C. Howe & Helen Nissenbaum, *Embodying Values in Technology: Theory and Practice*, in INFORMATION TECHNOLOGY AND MORAL PHILOSOPHY 322, 322 (Jeroen van den Hoven & John Weckert eds., 2008); Deirdre K. Mulligan & Helen Nissenbaum, *The Concept of Handoff as a Model for Ethical Analysis and Design*, in OXFORD HANDBOOK OF ETHICS OF AI 232, 233 (Markus D. Dubber, Frank Pasquale & Sunit Das eds., 2020); see generally MARY FLANAGAN & HELEN NISSENBAUM, *VALUES AT PLAY IN DIGITAL GAMES* (2014) (developing a framework for identifying socially recognized moral and political values in technology in the context of digital games).

151. Ahuja, *supra* note 144, at 1 (“OPM is pleased to issue skills-based hiring guidance and a competency model for Artificial Intelligence (AI), data, and technology talent to assist agencies to identify key skills and competencies needed for AI professionals and increase access to these technical roles for individuals with nontraditional academic backgrounds.”).

152. *DHS Generative AI Sector Playbook*, DEP'T OF HOMELAND SEC., <https://www.dhs.gov/ai> (last visited Aug. 21, 2024); Justin Doubleday, *DHS Sets 'Aggressive' Recruiting Strategy to Fill AI Jobs*, FED. NEWS NETWORK (Feb. 19, 2024), <https://federalnewsnetwork.com/artificial-intelligence/2024/02/dhs-sets-aggressive-recruiting-strategy-to-fill-ai-jobs/>.

Foreign Policy and National Security Workforce, Institutions, and Partnerships (NSM-3) and other presidential priorities.¹⁵³ The Department of Defense, with support from the OMB and the OSTP announced a new program which, if launched, would provide a reserve team of STEM and AI experts from academia and elsewhere that agencies could tap for short-term engagements to bring appropriate expertise in to assist with implementations, evaluations, and other work.¹⁵⁴ Centralizing the work of hiring and clearing these advisors and providing a simple mechanism for a range of agencies to bring them in for specific projects would expand the range of experts agencies could practically and financially afford to bring into their efforts.

Enforcement agencies with more dedicated technical expertise produced reports and held workshops to assist peer agencies; led calls to action clarifying the importance and need for technical experts; and created networks to build momentum domestically and internationally for increased technical expertise within consumer protection, competition, and civil rights enforcement agencies. For example, the FTC's Office of Technology issued a report designed to "establish a shared context and serve as a resource for building technical capacity in government agencies" and share information about how the Office of Technology "applies subject matter experts in regulatory and enforcement contexts."¹⁵⁵ The FTC took this effort to the international stage, initiating the International Competition Network Tech Forum's work to define best practices in building tech capacity in law enforcement agencies.¹⁵⁶ The Consumer Financial Protection Bureau's Office of Technology similarly sought to boost technical expertise across enforcement agencies.¹⁵⁷ They

153. *Fact Sheet: Biden-Harris Administration Announces Commitments from Across Technology Ecosystem Including Nearly \$100 Million to Advance Public Interest Technology*, WHITE HOUSE (July 16, 2024), <https://bidenwhitehouse.archives.gov/ostp/news-updates/2024/07/16/fact-sheet-biden-harris-administration-announces-commitments-from-across-technology-ecosystem-including-nearly-100-million-to-advance-public-interest-technology/>.

154. *Id.* (announcing the Trusted Advisors Pilot). This program was not operational by the end of the Biden-Harris Administration, and it is unclear whether work to stand it up is continuing or if the Trump Administration has abandoned it.

155. OFF. TECH. STAFF, FED. TRADE COMM'N, *BUILDING TECH CAPACITY IN LAW ENFORCEMENT AGENCIES* 3 (Mar. 2024), https://www.ftc.gov/system/files/ftc_gov/pdf/ot.techcapacityreport.pdf.

156. OFF. OF TECH., *Best Practices in Building Tech Capacity in Law Enforcement Agencies*, FED. TRADE COMM'N (Mar. 26, 2024), <https://www.ftc.gov/policy/advocacy-research/tech-at-ftc/2024/03/best-practices-building-tech-capacity-law-enforcement-agencies>; *see also Building Digital Capacity to Strengthen and Support Law Enforcement Agencies*, INT'L COMPETITION NETWORK, <https://www.internationalcompetitionnetwork.org/working-groups/icn-operations/technologists/technologist-forum-statement-on-building-agency-digital-capacity/> (last visited Sep. 19, 2025).

157. Erie Meyer, *Public Interest Tech Jobs: Regulate Tech and AI*, CONSUMER FIN. PROT. BUREAU (May 20, 2024), <https://www.consumerfinance.gov/about-us/blog/public-interest-tech-jobs-regulate-tech-and-ai/> (announcing the launch of a cross-government hiring effort

hosted trainings and briefings for enforcement agencies on numerous consumer financial protection topics, including emerging practices in biometrics and how to address algorithmic harms¹⁵⁸ and authored a guide to help enforcement agencies hire technologists to protect consumers.¹⁵⁹ In addition, EO 14110 directed the Civil Rights Division of the Department of Justice to convene federal civil rights offices to advance comprehensive use of their respective authorities and offices to prevent and address discrimination in the use of automated systems, including algorithmic discrimination, and “develop, as appropriate, additional training, technical assistance, guidance, or other resources . . . and consider providing [similar support] to State, local, Tribal, and territorial investigators and prosecutors.”¹⁶⁰

The White House and federal agencies took many other actions to build the staff of sociotechnical experts within agencies and directly available to them to support the responsible use and governance of AI.

2. *Building the Responsible AI Field*

The Administration further sought to build the field of responsible AI generally, rather than just the government capacity. The Administration established the U.S. AI Safety Institute (AIS), which is housed within NIST, to advance the science of AI safety; articulate, demonstrate, and disseminate the practices of AI safety; and support institutions, communities, and coordination around AI safety.¹⁶¹ To achieve these goals, the AIS conducts testing of advanced models and systems to assess potential and emerging risks; develops guidelines on evaluations and risk mitigations; and performs and

to embed technical experts across multiple agencies that share a variety of consumer protection, competition, and civil rights authorities).

158. Erie Meyer, *Bringing Tech Enforcers Together to Protect Consumers*, CONSUMER FIN. PROT. BUREAU (Mar. 14, 2023), <https://www.consumerfinance.gov/about-us/blog/bringing-tech-enforcers-together-to-protect-consumers/>.

159. *Hiring Technologists to Protect Consumers*, CONSUMER FIN. PROT. BUREAU, <https://www.consumerfinance.gov/about-us/careers/cfpb-technologist/hiring-technologists-to-protect-consumers/> (last modified Oct. 30, 2024).

160. Exec. Order No. 14,110 § 7.1(ii)–(iii), 88 Fed. Reg. 75191, 75211 (Nov. 1, 2023) (signed Oct. 30, 2023).

161. Press Release, Off. of Pub. Affs., at the Direction of President Biden, Department of Commerce to Establish U.S. Artificial Intelligence Safety Institute to Lead Efforts on AI Safety (Nov. 1, 2023), <https://www.commerce.gov/news/press-releases/2023/11/direction-president-biden-department-commerce-establish-us-artificial>. The Trump Administration renamed it to the Center for AI Standards and Innovation and narrowed its agenda. Press Release, Off. of Pub. Affs., Statement from U.S. Secretary of Commerce Howard Lutnick on Transforming the U.S. AI Safety Institute into the Pro-Innovation, Pro-Science U.S. Center for AI Standards and Innovation (June 3, 2025), <https://www.commerce.gov/news/press-releases/2025/06/statement-us-secretary-commerce-howard-lutnick-transforming-us-ai>.

coordinates technical research.¹⁶² To enable this work, AISI works closely with experts from across the AI industry, civil society, and sister safety institutes.¹⁶³

AISI and the broader set of NIST experts working on AI were tasked with providing technical guidance to support the responsible development and use of AI. As noted above, AISI and NIST issued guidance—some still in progress—on a range of technical issues, including guidance for AI developers in managing the evaluation of misuse of dual-use foundation models, frameworks on managing generative AI risks and securely developing generative AI systems and dual-use foundation models, and provided a technical report to the White House outlining tools and techniques to reduce the risks from synthetic content.¹⁶⁴ All of these documents were produced with input from stakeholders.

To support stakeholder participation in AI governance, AISI established a consortium of over 200 AI stakeholders that seeks to “unite AI creators and users, academics, government and industry researchers, and civil society organizations in support of the development and deployment of safe and trustworthy artificial intelligence.”¹⁶⁵ The consortium includes a wide range of stakeholders including AI companies like Anthropic and OpenAI, technology companies like Apple and Google, energy companies including PG&E, chip manufacturers like NVIDIA, and universities including NYU, Syracuse, and UC Berkeley.¹⁶⁶ The consortium has working groups on topics including Risk Management for Generative AI, Synthetic Content, Capability Evaluations, Red-Teaming, and Safety & Security.¹⁶⁷

AISI began to address the access challenges that stymie stakeholders’ full participation in decisions about AI models. AISI signed memorandums of understanding (MOU) with two major AI companies, Anthropic and OpenAI, that “enable formal collaboration on AI safety research, testing and

162. See generally U.S. A.I. SAFETY INST., NAT’L INST. OF STANDARDS & TECH., *supra* note 123.

163. *Id.* at 2.

164. See, e.g., *Department of Commerce Announces New Guidance, Tools 270 Days Following President Biden’s Executive Order on AI*, NAT’L INST. OF STANDARDS & TECH. (July 26, 2024), <https://www.nist.gov/news-events/news/2024/07/departments-commerce-announces-new-guidance-tools-270-days-following>.

165. *Biden-Harris Administration Announces First-Ever Consortium Dedicated to AI Safety*, NAT’L INST. OF STANDARDS & TECH. (Feb. 8, 2024), <https://www.nist.gov/news-events/news/2024/02/biden-harris-administration-announces-first-ever-consortium-dedicated-ai>.

166. *Artificial Intelligence Safety Institute Consortium: AISIC Members*, NAT’L INST. OF STANDARDS & TECH., <https://www.nist.gov/artificial-intelligence/artificial-intelligence-safety-institute-consortium-aisic/aisic-members> (last visited Sep. 17, 2025).

167. *Artificial Intelligence Safety Institute Consortium: AISIC Working Groups*, NAT’L INST. OF STANDARDS & TECH., <https://www.nist.gov/artificial-intelligence/artificial-intelligence-safety-institute-consortium-aisic/aisic-working> (last visited Sep. 17, 2025).

evaluation” between these companies and the federal government.¹⁶⁸ Each company’s MOU “establishes the framework for the U.S. AI Safety Institute to receive access to major new models from each company prior to and following their public release” to enable collaborative research on AI model capabilities, safety risks, and risk mitigation.¹⁶⁹

To advance global coordination on AI governance, AISI hosted the inaugural convening of an International Network of AI Safety Institutes in San Francisco. Secretary of Commerce Gina Raimondo welcomed delegations from AI Safety Institutes in Australia, Canada, the European Union, France, Japan, Kenya, the Republic of Korea, Singapore, and the United Kingdom.¹⁷⁰ The goal of the convening was to “kickstart the Network’s technical collaboration ahead of the AI Action Summit in Paris in February 2025.”¹⁷¹ The convening also included “experts from international civil society, academia, and industry” who would “inform the work of the Network and ensure a robust view of the latest developments in the field of AI.”¹⁷²

In the leadup to the convening, AISI formed the Testing Risks of AI for National Security (TRAINS) Taskforce, which is chaired by AISI and includes representation from the Department of Defense, including the Chief Digital and Artificial Intelligence Office and the National Security Agency; the Department of Energy and ten of its National Laboratories; the Department of Homeland Security, including the Cybersecurity and Infrastructure Security Agency; and the National Institutes of Health at the Department of Health and Human Services.¹⁷³ TRAINS aimed to enable coordinated research and testing of advanced AI models on issues related to national security, including “radiological and nuclear security, chemical and biological security, cybersecurity, critical infrastructure, [and] conventional military capabilities.”¹⁷⁴

168. *U.S. AI Safety Institute Signs Agreements Regarding AI Safety Research, Testing and Evaluation with Anthropic and OpenAI*, NAT’L INST. OF STANDARDS & TECH. (Aug. 29, 2024), <https://www.nist.gov/news-events/news/2024/08/us-ai-safety-institute-signs-agreements-regarding-ai-safety-research>.

169. *Id.*

170. Press Release, U.S. Dep’t of Com., U.S. Secretary of Commerce Raimondo and U.S. Secretary of State Blinken Announce Inaugural Convening of International Network of AI Safety Institutes in San Francisco (Sep. 18, 2024), <https://www.commerce.gov/news/press-releases/2024/09/us-secretary-commerce-raimondo-and-us-secretary-state-blinken-announce>.

171. *Id.*

172. *Id.*

173. Press Release, U.S. Dep’t of Com., 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.commerce.gov/news/press-releases/2024/11/us-ai-safety-institute-establishes-new-us-government-taskforce>.

174. *Id.*

The Biden-Harris Administration invested in building the AI research ecosystem necessary to support the use of AI for important public missions and to establish a strong, sociotechnical understanding of risks and the methods and tools to address them. It created eighteen new AI institutes across the United States through the National Science Foundation-led National Artificial Intelligence Research Institutes program—a research investment that began in August of 2020 during the Trump-Pence Administration which established the first seven national AI research institutes.¹⁷⁵ Many of these institutes enjoy private support as well, but the public investment ensures they are directed towards research that will benefit the public, ranging from promoting ethical and trustworthy AI systems and technologies, developing novel approaches to cybersecurity, addressing climate change, expanding our understanding of the brain, and enhancing education and public health.¹⁷⁶ In addition, it launched the National Science Foundation’s (NSF) Responsible Design, Development, and Deployment of Technologies (ReDDDoT) program that supports multidisciplinary, multi-sector teams that examine and demonstrate the principles, methodologies, implementations, and impacts associated with responsible design, development, and deployment of technologies in practice.¹⁷⁷ A collaboration between the NSF and philanthropic funders, the program and other publicly funded efforts support research that is developing new methods and approaches to ensure that ethical, legal, and societal considerations and community values are embedded across technology lifecycles to generate products that promote the public’s well-being and mitigate harm.

The Administration launched the National AI Research Resource (NAIRR) pilot—a national infrastructure led by the NSF in partnership with the Department of Energy and other governmental and nongovernmental partners—that makes available resources to support the nation’s AI research and education community.¹⁷⁸ The NAIRR supports research teams across forty-nine states that are tackling projects covering deepfake detection, AI safety, next-generation medical diagnoses, environmental protection, and

175. Michael Kratsios & Chris Liddell, Off. Sci. & Tech. Pol’y, *The Trump Administration Is Investing \$1 Billion in Research Institutes to Advance Industries of the Future*, WHITE HOUSE (Aug. 26, 2020), <https://trumpwhitehouse.archives.gov/articles/trump-administration-investing-1-billion-research-institutes-advance-industries-future/>.

176. *NSF Announces \$100 Million Investment in National Artificial Intelligence Research Institutes Awards to Secure American Leadership in AI*, NAT’L SCI. FOUND. (July 29, 2025), <https://www.nsf.gov/news/nsf-announces-100-million-investment-national-artificial> (discussing public-private partnerships).

177. *Responsible Design, Development, and Deployment of Technologies*, NAT’L SCI. FOUND. (Jan. 8, 2024), <https://www.nsf.gov/funding/opportunities/redddot-responsible-design-development-deployment-technologies/506215/nsf24-524>.

178. NAT’L A.I. RSCH. RES. PILOT, <https://nairrpilot.org> (last visited Sep. 17, 2025).

materials engineering.¹⁷⁹ It provides public infrastructure to support research necessary to address AI governance challenges.¹⁸⁰

These public research investments are essential for progress on AI Governance. For example, the Defense Advanced Research Projects Agency (DARPA) launched one of the first and most recognized programs in the area of Explainable AI (XAI) with the goal of enabling end users to better understand, trust, and effectively manage artificially intelligent systems.¹⁸¹ This early public research investment sparked broad interest in an area essential to AI governance.¹⁸² And this is just one example. A robust publicly funded research ecosystem will produce the methods and tools to ensure AI systems are fit for purpose, trustworthy, rights-respecting, and safe. It will also ensure AI research advances our nation's grand ambitions and addresses our gravest risks, whether they arise from adversarial nations seeking to undermine our national security, or decades of inaction to address the looming climate crisis.

Together these actions bolstered both the technical workforce and technical expertise in government, as well as in civil society, academia, and other stakeholders. The Administration invested in the research and education to create AI governance methods and future practitioners, galvanized support for the field of public interest technology, provided resources to support AI research and applications outside large companies, and created momentum for similar international efforts. The Administration created public venues for all stakeholders to participate in “governance-by-design” work. These efforts help ensure AI governance is consistent with the norms of public governance and designed to center and be responsive to public values and not just private interests.

179. To see the range of projects supported, see *Resource Allocation*, NAT'L A.I. RSCH. RES. PILOT, <https://nairrpilot.org/projects/awarded> (last visited Sep. 17, 2025).

180. *Democratizing the Future of AI R&D: NSF to Launch National AI Research Resource Pilot*, NAT'L SCI. FOUND. (Jan. 24, 2024), <https://www.nsf.gov/news/democratizing-future-ai-rd-nsf-launch-national-ai-research>.

181. *XAI: Explainable Artificial Intelligence*, DEF. ADVANCED RSCH. PROJECTS AGENCY, <https://www.darpa.mil/research/programs/explainable-artificial-intelligence> (last visited Sep. 17, 2025).

182. David Gunning, Eric Vorm, Jennifer Yunyan Wang & Matt Turek, *DARPA's Explainable AI (XAI) Program: A Retrospective*, 2 APPLIED AI LETTERS e61 (2021) (claiming the program stimulated the field of explainable AI research and “produced a more nuanced understanding of XAI uses and users, the psychology of XAI, the challenges of measuring explanation effectiveness, as well as producing a new portfolio of XAI ML and HCI techniques”); Atul Rawal, James McCoy, Danda B. Rawat, Brian M. Sadler & Robert St. Amant, *Recent Advances in Trustworthy Explainable Artificial Intelligence: Status, Challenges, and Perspectives*, 3 IEEE TRANSACTIONS ON A.I. 852, 852–53 (2022) (Figure 1 showing the rise in explainable AI research papers in 2017 and noting its “emergence along with the U.S. DoD DARPA XAI program”).

C. MAINTAINING THE PUBLICNESS OF POLICYMAKING: FOCUSING ON IMPACT RATHER THAN SYSTEMS AND REQUIRING STAKEHOLDER PARTICIPATION THROUGHOUT THE AI LIFECYCLE

The Administration sought to maintain the publicness of the values and policies embedded in AI systems and provide robust opportunities for public input into its deliberations to intervene in technological design. It did so in general and specific ways.

This Section first discusses two interventions that reframed the project of AI governance away from models and systems towards impacts on individuals' and communities' rights and safety: the Blueprint for an AI Bill of Rights¹⁸³ and the turn to *use cases* as the focus of evaluation in federal AI governance. This Section next describes in detail how this reframing along with new system documentation and public engagement requirements in the OMB guidance to federal agencies on the development and use of AI maintained the visibility and attention to rights and values during agency design and deployment practices. Finally, this Section documents the Administration's approach to a significant policy—the availability of model-weights for powerful AI models—which engaged the public in considerations of the breadth of values implicated in a governance-by-design strategy and produced a policy that exhibits modesty and restraint in design.

1. *Reframing The Project of AI Governance*

Two overarching and underappreciated interventions reoriented the debates about AI governance: the Blueprint for the AI Bill of Rights¹⁸⁴ and the use-case orientation of the federal guidance on the development and use of AI.

a) The AI Bill of Rights

First, as Alondra Nelson, former Principal Deputy Director for Science and Society at the White House OSTP, Deputy Assistant to the President, who led the creation of the White House Blueprint for an AI Bill of Rights (AI BoR), explained, the AI BoR acts as “civic architecture” “creating infrastructure for collective participation in AI policy.”¹⁸⁵ It grounded AI governance in the achievement of rights and civil liberties, establishing “a sociotechnical approach that recalibrates the relationship between technology

183. AI BoR, *supra* note 87.

184. *Id.*

185. Alondra Nelson, *From Blueprint to Building Blocks: The AI Bill of Rights as Civic Architecture 2* (forthcoming) (on file with authors); *see also* Alondra Nelson, Inst. for Advanced Study, Presentation at Artificial Intelligence and Democratic Freedoms, Symposium by Knight First Amendment Institute at Columbia University (Apr. 10, 2025), <https://knightcolumbia.org/events/artificial-intelligence-and-democratic-freedoms>.

and society by positioning individuals not merely as passive users of AI systems but as rights-bearing individuals and communities with legitimate claims to protection, agency, and redress.”¹⁸⁶ The AI BoR did more than establish guiding principles, it established a frame and reference point for the public to see themselves, their communities, and the stakes in what were technocratic, expert-dominated debates.

Second, and perhaps stealthier to those outside the government, the OMB guidance to agencies on the responsible development and use of AI focused on the evaluation of AI *use cases*. Typically, the object of assessment or evaluation is a *system*¹⁸⁷ generally exclusive to the technical aspects, not the institutional aspects, that together co-create a system’s ultimate impact. By reorienting the government’s analysis of AI around *use cases*, such as use of an AI to assist in benefits determinations, the Administration established a new paradigm for accounting for the potential outcomes of incorporating AI into a government process. This reorientation adopts a sociotechnical approach to risk management found in high-risk fields.¹⁸⁸

Focusing on use cases addresses shortcomings that frequently pervade other methods of assessment. As technology scholar Roel Dobbe has explained, safety science research reveals that “systems cannot be safeguarded by technical design choices on the model or algorithm alone.”¹⁸⁹ Rather, he explains, it is necessary to take an “end-to-end” approach to analyzing risks and a sociotechnical system—a perspective that considers “the context of use, impacted stakeholders . . . and informal institutional environment” when deploying mitigations.¹⁹⁰

186. Nelson, *From Blueprint to Building Blocks*, *supra* note 185, at 1.

187. For example, 44 U.S.C. § 3501 note (2000 & 2002 Amendments), and OFF. OF MGMT. & BUDGET, EXEC. OFF. OF THE PRESIDENT, MEMO. NO. M-03-22, OMB GUIDANCE FOR IMPLEMENTING THE PRIVACY PROVISIONS OF THE E-GOVERNMENT ACT OF 2002 (Sep. 26, 2003) (requiring agencies to conduct a Privacy Impact Assessment before “developing or procuring information technology that collects, maintains, or disseminates information that is in an identifiable form”).

188. One notable exception to the use-case orientation is the Administration’s focus on the risks posed by dual-use foundation models with widely available model weights, discussed, *infra*, in Part IV(C)(2), which takes a systems analysis. However, as explained, *infra*, the inquiry led by the National Telecommunications and Information Administration (NTIA) explored contextualized risks and benefits of systems in various domains of use, again centering expected impact rather than abstract risk.

189. Roel I. J. Dobbe, *System Safety and Artificial Intelligence*, in THE OXFORD HANDBOOK OF AI GOVERNANCE 441, 441 (Justin B. Bullock et al., eds., 2022).

190. *Id.*

By considering policies, users, the technical interfaces, and outputs in context, a use-case orientation thus avoids “traps”¹⁹¹ researchers have identified in detached evaluation of technical systems alone. Specifically, such a sociotechnical system framing helps resist abstractions common in technical practice that push important normative decisions out of view,¹⁹² rendering visible built-in politics and values, and opening them up for contestation.¹⁹³

Such visibility and contestation is necessary to ensure a systemic orientation towards public values, and away from what Cohen and Waldman have called regulatory managerialism—the importing of the “practices for organizing and overseeing private sector, capitalist economic production and . . . the logics and underlying ideologies in which those practices are rooted” into regulated activities.¹⁹⁴ As scholar of regulation Christie Ford has powerfully argued, regulatory managerialism contributes to “peoples’ . . . alienation from public institutions, and the perspectives of the regulators who are supposed to be safeguarding their interests.”¹⁹⁵

The focus on *use cases* thus responds to the limits of algorithmic accountability,¹⁹⁶ and calls to introduce more qualitative elements into assessments,¹⁹⁷ including scenario analysis, which produces an “extended narrative prediction of how a given policy decision will increase the likelihood of some complex set of consequences and decrease that of others.”¹⁹⁸ Such tools “enable policy evaluators to better understand and predict interrelated aspects of technological advance, economic changes, and policy shifts,”¹⁹⁹

191. Andrew D. Selbst, Danah Boyd, Sorelle A. Friedler, Suresh Venkatasubramanian & Janet Vertesi, *Fairness and Abstraction in Sociotechnical Systems*, FAT* ’19: PROCS. OF THE CONF. ON FAIRNESS, ACCOUNTABILITY, & TRANSPARENCY 59 (2019).

192. *Id.*

193. See Mulligan & Bamberger, *Procurement as Policy*, *supra* note 1, at 842–57 (discussing the importance of political-visibility-enhancing processes and the resulting contestability).

194. Cohen & Waldman, *supra* note 26, at iv.

195. Christie Ford, *Regulation as Respect*, 86 L. & CONTEMP. PROBS. 133, 134 (2023).

196. See *id.* at 138 (“[Regulators] tend not even to have methods for determining which specific groups of people—especially vulnerable ones, however defined—should be considered as part of any regulatory impact assessment. Nor do they have decision rules for how to determine the balances between benefits and costs when considering disaggregated groups. These gaps highlight the level of abstraction from real humans at which managerialism operates, and they make inequities harder to see.”).

197. Andrew D. Selbst, *An Institutional View of Algorithmic Impact Assessments*, 35 HARV. J. L. & TECH. 117, 122, 180 (2021); see also Rory Van Loo, *Stress Testing Governance*, 75 VAND. L. REV. 553, 558 (2022) (advocating the use of “stress tests” that “incorporate elements of scenario analysis and simulations,” which “go beyond modeling threats . . . to assess how well private or public organizations would respond to those threats”); Frank Pasquale & Glyn Cashwell, *Four Futures of Legal Automation*, 63 UCLA L. REV. DISCOURSE 26, 28 (2015) (arguing for a scenario analysis in a “time of rapid technological change”).

198. Pasquale, *supra* note 26, at 56.

199. *Id.*

rather than “doubling down on the managerialist impulse to quantify costs and benefits.”²⁰⁰

Orienting AI assessment towards use cases centers the analysis on the distributional and other implications of actual applications on real sets of human beings in concrete contexts. It rejects assessing risk based on what Science and Technology Studies (STS) scholar Donna Haraway calls the “gaze from nowhere.”²⁰¹ It contests what Professor Sheila Jasanoff argues is risk regulation’s “favored . . . type of objectivity”; specifically, it challenges “claims making [that] achieves power by ostensibly detaching knowledge from potentially biased standpoints and from the distortions that any perspective or viewpoint necessarily entails . . . the kind of purification that scientists have historically aimed for in making representations of nature.”²⁰² Instead, a focus on use cases situates the understanding of a system’s efficacy, and its risks, in the messy entangled world in which it is located, stabilized, and iteratively repaired, so it can work.

Ideally, then, adopting a use-case orientation moves designers, data scientists, software engineers, and the myriad of others involved in designing, using, justifying, and repairing systems towards Professor Lucy Suchman’s concept of “located accountabilities,” which replaces “ways of being nowhere while claiming to see comprehensively” with “views from somewhere.”²⁰³ These are views built on “partial, locatable, critical knowledges,”²⁰⁴ and demand that designers take responsibility for what they see and what they “learn how to build.”²⁰⁵

In sum, the shift in analysis from systems to *use cases* reframes the stakes of AI governance towards a focus on individuals’ rights and safety. Stakeholder engagement draws designers’ and users’ attention to the rights of and impacts on relevant populations. Requiring the involvement of internal subject matter experts responsible for attending to rights within government brings privacy, accessibility, civil rights and security into processes that can shape the design and deployment of systems. Together, these practices center the context in which AI is just one component and directs agencies to question how its introduction will affect agencies’ missions and impact the public they serve. By injecting the public’s voice and internal rights-oriented expertise into

200. *Id.*

201. Donna Haraway, *Situated Knowledges: The Science Question in Feminism and the Privilege of Partial Perspective*, 14 FEMINIST STUD. 575, 581 (1988).

202. Sheila Jasanoff, *The Practices of Objectivity in Regulatory Science*, in SOCIAL KNOWLEDGE IN THE MAKING 307, 309 (Charles Camic, Neil Gross & Michèle Lamont eds., 2011).

203. Lucy Suchman, *Located Accountabilities in Technology Production*, 14 SCANDINAVIAN J. INFO. SYS. 91, 94 (2002) (quoting Donna Haraway, *supra* note 201, at 590).

204. *Id.* at 96 (quoting Donna Haraway, *supra* note 201, at 584).

205. *Id.*

technological design choices the Administration sought to make AI use responsive to public needs and consistent with democratic values.

b) OMB Guidance to Federal Agencies

As described, *supra*, the White House OMB issued guidance to federal agencies on the responsible development and use of AI, and accompanying procurement guidance.²⁰⁶ The OMB guidance requires agencies to engage affected stakeholders in the design, including risk mitigation choices, of AI systems used by agencies.²⁰⁷ The federal guidance reflected and built on the Administration's commitments both to public participation and community engagement²⁰⁸ to developing a more effective set of strategies and tools to support meaningful public participation and community engagement in government policy-making and service design and delivery. At a time of growing distrust in institutions and disaffection with government, the Administration considered these efforts essential to building public trust in government; designing effective, inclusive and accessible policies and services to serve the full public; and aligning government practice with democratic ideals of a government of the people, by the people, and for the people.²⁰⁹ In addition, it reflected the Biden-Harris Administration's stated commitment to advancing equity across services, policies, and programs, a commitment reflected in the two equity executive orders²¹⁰ as well as the effort at modernizing regulatory review to account for the distributional consequences of regulation and to ensure that regulatory initiatives do not unduly burden the disadvantaged.²¹¹

The guidance established a more demanding set of risk management processes and requirements for government's use of rights-impacting AI, including facial recognition technologies. It centered the engagement of key government experts with responsibility for rights and consultation with and input from affected communities and the public throughout the AI lifecycle.²¹²

206. OMB MEMO. M-24-10, *supra* note 101; OMB MEMO. M-24-18, *supra* note 105.

207. OMB MEMO. M-24-10, § 5(v)(B), *supra* note 101, at 22 ("Consult and incorporate feedback from affected communities and the public.").

208. OFF. OF MGMT. & BUDGET, EXEC. OFF. OF THE PRESIDENT, MEMO. NO. M-25-07, BROADENING PUBLIC PARTICIPATION AND COMMUNITY ENGAGEMENT IN THE REGULATORY PROCESS (Jan. 15, 2025), <https://bidenwhitehouse.archives.gov/wp-content/uploads/2025/01/M-25-07-Broadening-Participation-and-Engagement.pdf>.

209. Methods and Leading Practices for Advancing Public Participation and Community Engagement with the Federal Government, 89 Fed. Reg. 19885, 19886 (Mar. 20, 2024).

210. Exec. Order No. 13,985, 86 Fed. Reg. 7009 (Jan. 20, 2021); Exec. Order No. 14,091, 88 Fed. Reg. 10825 (Feb. 16, 2023).

211. Memorandum on Modernizing Regulatory Review, 86 Fed. Reg. 7223 (Jan. 26, 2021).

212. Deirdre K. Mulligan, Principal Deputy U.S. Chief Tech. Officer, White House Off. of Sci. & Tech. Pol'y, Testimony on the Civil Rights Implications of the Federal Use of

And it further required federal agencies to make the goals, policies, and values embedded in systems and animating their use open and subject to public input.²¹³

An interrelated set of requirements build visibility into the AI lifecycle including:

- Data documentation exposing the provenance, the quality and representativeness in relation to purpose, an assessment of its breadth and gaps and how shortcomings of the data have been addressed by the agency or vendor, and if the data is maintained by the Federal Government, whether that it is publicly disclosable as an open government data asset;²¹⁴
- A public use case inventory with accessible documentation in plain language of the system's functionality to serve as public notice of the AI to its users and the general public;²¹⁵
- A requirement for reasonable and timely notice about the use of the AI to those subject to them and a means to directly access any public documentation about it in the use case inventory;²¹⁶
- Notice to individuals when the use of the AI results in an adverse decision or action that specifically concerns them, explanations for such decisions and actions, and if applicable their right to appeal;²¹⁷ and,
- A required human fallback and escalation system so impacted individuals can appeal or contest an AI use case's negative impacts²¹⁸ and mechanisms for individuals to choose a human alternative where practicable and consistent with law.²¹⁹

A second set of interlocking requirements creates assessments of algorithmic processes that bridge between technical and agency experts and the publics and communities they serve, as well as outside experts:

- Consultation with affected communities, including underserved communities on the design, development, and use of the AI and risk mitigations;²²⁰

Facial Recognition Technology Before the U.S. Commission on Civil Rights (Mar. 8, 2024), <https://www.usccr.gov/files/2024-04/frt-transcript.pdf>.

213. OMB MEMO. M-24-10, *supra* note 101.

214. *Id.* §§ 4(d)(ii), 5(c)(iv)(A)(3).

215. *Id.* § 3(a)(iv).

216. *Id.* § 5(c)(iv)(I).

217. *Id.* § 5(c)(v)(D).

218. *Id.* § 5(c)(v)(E).

219. *Id.* § 5(c)(v)(F).

220. *Id.* § 5(c)(v)(B).

- Impact assessments to document the intended purpose for the AI and its expected benefit, potential risks, and quality of the relevant data;²²¹
- Testing requirements for performance in real-world contexts;²²²
- Requirements for agencies to identify, assess, and mitigate algorithmic discrimination and harmful bias to ensure that federal government use of AI does not decrease equity or fairness;²²³ and,
- Ongoing monitoring and thresholds for periodic human review.²²⁴

These complement existing requirements such as privacy impact assessments²²⁵ that are intended to bridge the gap between internal experts and the outside world.

Together, these requirements make the values and policy choices built into system designs visible. For those designing and deploying systems, these provide scaffolding for what Phil Agre called “critical technical practice,”²²⁶ in contrast to the formalistic and universalizing tendencies of the field of AI. Agre called for methods and practices of studying, building, and implementing technology that straddled the “craft work of design”²²⁷ and the “reflexive work of critique”²²⁸ to reveal the value choices inherent in technical terminology and design and move towards the useful context specific implementations required for AI to be meaningfully useful. They reflect work in the field of responsible

221. *Id.* § 5(c)(iv)(A).

222. *Id.* § 5(c)(iv)(B).

223. *Id.* § 5(c)(v)(A).

224. *Id.* § 5(c)(iv)(E).

225. DEPT OF HOMELAND SEC., PRIVACY IMPACT ASSESSMENTS: THE PRIVACY OFFICE OFFICIAL GUIDANCE (June 2010), https://www.dhs.gov/xlibrary/assets/privacy/privacy_pia_guidance_june2010.pdf.

226. Philip Agre, *Toward a Critical Technical Practice: Lessons Learned in Trying to Reform AI*, in SOCIAL SCIENCE, TECHNICAL SYSTEMS, AND COOPERATIVE WORK: BEYOND THE GREAT DIVIDE 155 (Geoffrey Bowker, Susan Leigh Star, Les Gasser & William Turner eds., 1997).

227. *Id.*

228. *Id.*

AI that has advanced data and model documentation,²²⁹ auditing,²³⁰ and impact assessments as well as participatory and co-design.²³¹

For impacted communities, the system documentation along with requirements for real-world testing and impacted community involvement in risk-mitigation choices, not just design, provide unprecedented opportunities to shape the use of technology. With the use-case orientation, these requirements keep the focus on designing within context and assessing how AI systems affect the outcomes for real people. The notices to negatively impacted individuals and requirements for human fallback and redress create ongoing moments of visibility making space for “technological dramas”²³²—contestation and reexamination of the policies and values baked into the sociotechnical system. These processes acknowledge the inevitability of “algorithmic breakdowns.”²³³ They create intentional seams that expose the configurability, complexity, and fragility of systems and actively resist the tendency for technological infrastructure to recede into the background.²³⁴

Finally, Algorithmic Impact Assessments (AIAs) are required to examine the efficacy and risks, including potential bias, in AI systems.²³⁵ They are

229. Margaret Mitchell, Simone Wu, Andrew Zaldivar, Parker Barnes, Lucy Vasserman, Ben Hutchinson, Elena Spitzer, Inioluwa Deborah Raji & Timnit Gebru, *Model Cards for Model Reporting*, FAT* 19: PROCS. OF THE CONF. ON FAIRNESS, ACCOUNTABILITY, & TRANSPARENCY 220 (2019); Ben Hutchinson, Andrew Smart, Alex Hanna, Emily Denton, Christina Greer, Oddur Kjartansson, Parker Barnes & Margaret Mitchell, *Towards Accountability for Machine Learning Datasets: Practices from Software Engineering and Infrastructure*, FACCT '21: PROCS. OF THE 2021 ACM CONF. ON FAIRNESS, ACCOUNTABILITY, & TRANSPARENCY 560, 560–75 (2021); Mark P. Sendak, Michael Gao, Nathan Brajer & Suresh Balu, *Presenting Machine Learning Model Information to Clinical End Users with Model Facts Labels*, 3 NPJ DIGIT. MED. 41 (2020).

230. Inioluwa Deborah Raji, Andrew Smart, Rebecca N. White, Margaret Mitchell, Timnit Gebru, Ben Hutchinson, Jamila Smith-Loud, Daniel Theron & Parker Barnes, *Closing the AI Accountability Gap: Defining an End-To-End Framework for Internal Algorithmic Auditing*, FAT*20: PROCS. OF THE 2020 CONF. ON FAIRNESS, ACCOUNTABILITY, & TRANSPARENCY 33–44 (2020).

231. For a recent overview, see Ned Cooper, Tiffanie Horne, Gillian R. Hayes, Courtney Heldreth, Michal Lahav, Jess Holbrook & Lauren Wilcox, *A Systematic Review and Thematic Analysis of Community-Collaborative Approaches to Computing Research*, CHI '22: PROCS. OF THE 2022 CHI CONF. ON HUMAN FACTORS IN COMPUTING SYS. 1–18 (2022); see also SASHA COSTANZA-CHOCK, *DESIGN JUSTICE: COMMUNITY-LED PRACTICES TO BUILD THE WORLDS WE NEED* (2020).

232. Bryan Pfaffenberger, *Technological Dramas*, 17 SCI., TECH., & HUM. VALUES 282 (1992).

233. Deirdre K. Mulligan & Daniel S. Griffin, *Rescripting Search to Respect the Right to Truth*, 2 GEO. L. TECH. REV. 557, 559 (2018).

234. For a review of the concept of “seamful design,” see Sarah Inman & David Ribes, *“Beautiful Seams”: Strategic Revelations and Concealments*, CHI '19: PROCS. OF THE 2019 CHI CONF. ON HUMAN FACTORS IN COMPUTING SYS. (2019).

235. OMB MEMO. M-24-10 § 5(c)(iv)(A), (c)(v)(A), *supra* note 101.

intended to ensure that AI systems used by federal agencies do not reproduce existing patterns of discrimination, inhere the prejudice of prior decision makers, or simply reflect the widespread biases that persist in society. Ideally moving away from the litany of existing algorithmic decision-making systems that have been shown to “automate inequality”²³⁶ perpetuating biases of various forms in high-stake consequences including the termination of welfare benefits, granting or denying immigration visas or wrongful imprisonment based on biased AI facial recognition.

Impact assessments, as legal scholar Andrew D. Selbst suggests, encourage those who build systems to think critically about the “potential impacts of a complex project before its implementation, thereby heading off risks before they become too costly to correct.”²³⁷ They further create documentation of decisions made during AI systems development to promote accountability for those decisions, as well as to provide useful information for policy interventions to correct for bias down the line.²³⁸ For these reasons, Danielle Keats Citron and Frank Pasquale have advocated for the use of privacy and civil liberties impact assessments when evaluating an AI scoring system’s “negative, disparate impact on protected groups, arbitrary results, mischaracterizations, and privacy harms.”²³⁹

At their best, AIAs promote reflexivity, encouraging decisionmakers to document and reflect on their assumptions, their data, and their measurement models at the earliest stages of the development process. Ideally, they act as a *boundary object* offering a shared reference point that enables broader participation in efforts to shape the role of these increasingly consequential technologies in society.

Julie Cohen and Ari Ezra Waldman caution, however, AIAs, as well as other impact or risk assessment frameworks can—like other “regulatory managerial” practices—allow entities to conceal predatory behavior behind a facade of procedural legitimacy.²⁴⁰ They can also marginalize outside or expert perspectives on the relevant interests at stake (e.g., privacy, racial equity) in favor of “check-box compliance sensibilities.”²⁴¹ In our own work, therefore, we have foregrounded the potential benefits of using more focused assessment techniques, such as human rights impact assessments (HRIAs), to address

236. See generally VIRGINIA EUBANKS, *AUTOMATING INEQUALITY: HOW HIGH-TECH TOOLS PROFILE, POLICE, AND PUNISH THE POOR* (2018).

237. Selbst, *supra* note 197, at 122.

238. *Id.* at 118.

239. Danielle Keats Citron & Frank A. Pasquale, *The Scored Society: Due Process for Automated Predictions*, 89 WASH. L. REV. 1, 26 (2014).

240. Cohen & Waldman, *supra* note 26, at v.

241. *Id.* at vi.

public values more systematically in technical practice.²⁴² HRIAs focus on impacts on affected communities and center rights.²⁴³ These attributes pull towards a situated consideration of impacts rather than an abstract conversation about system properties. The OMB use-case orientation for AIA similarly foregrounds the institutional and social, along with the technical, and requires agencies to assess how AI systems understood as reconfigurations of work practices or functions²⁴⁴ will impact the rights and safety of people.

Seen in this light, the AI BoR and the *use-case* framing together reoriented the project and stakes of AI governance and created both a “civic architecture” to support public participation and demanding responsibility and accountability for the impact of technology design and use. Operating within these frames, the OMB requirements provided scaffolding for “critical technical practice”²⁴⁵ and for ongoing public input and contestation about systems values and impacts throughout the lifecycle.

2. *The National Telecommunications and Information Administration (NTIA) and Model Weights*

The NTIA public process and report on large AI models with widely available weights—the numerical parameters that comprise an AI model determine its behavior—created a venue to foster the meta-discussion about when and whether it is appropriate to direct aspects of the design of AI models in the service of specific values at all and if so, how to prioritize among values. These are exactly the cross-cutting discussions about whether to engage in a “Governance-by-Design” initiative that we found missing, and without a clear home, in our prior work.²⁴⁶ The public nature of this process underscored the role government could play in determining whether models with a key property—disclosed weights—were available in the domestic market. The publicness of the process shed light on the possibility that the government would directly shape technology to effect its goal of reducing the risks of AI. The publicness of the process increased the chance that if the

242. Bamberger & Mulligan, *Governance-By-Design*, *supra* note 2, at 764.

243. See THE DANISH INST. FOR HUMAN RIGHTS, HUMAN RIGHTS IMPACT ASSESSMENT GUIDANCE AND TOOLBOX 8 (2020), https://www.humanrights.dk/files/media/document/DIHR%20HRIA%20Toolbox_Welcome_and_Introduction_ENG_2020.pdf (“Engagement with rights-holders and other stakeholders is essential . . . requiring background research and fieldwork, as well as heavily based on the participation of rights-holders other stakeholders”).

244. Mulligan & Nissenbaum, *supra* note 150, at 233 (explaining the ways in which reconfiguring how a function is performed can impact rights and values).

245. Kirsten Boehner, Shay David, Joseph ‘Jofish’ Kaye & Phoebe Sengers, *Critical Technical Practice as a Methodology for Values in Design*, in CHI 2005 WORKSHOP: QUALITY, VALUES AND CHOICES, at 1 (“Critical Technical Practice (CTP) is an approach to identifying and altering philosophical assumptions underlying technical practice.”).

246. Bamberger & Mulligan, *Governance-By-Design*, *supra* note 2, at 759–70.

government decided to limit the availability of models with widely disclosed weights it would be understood as a modality of regulation rather than a technological inevitability or market decision. It created the conditions for government to be held responsible for the political consequences of shaping technology.²⁴⁷

In *Saving Governance-by-Design*, we explained that “[e]xisting governance institutions often lack . . . [the] substantive regulatory capacity—breadth of authority, competence, and vision . . . to support the rational use of technology to govern.”²⁴⁸ We then suggested a number of approaches, including “coordination and input from a range of government actors; and [c]onditioning governance-by-design on multi-stakeholder involvement” that could “broaden the set of values that decision makers must consider, [and] decision makers’ capacity to address relevant values.”²⁴⁹ The NTIA process did both.

The President, through the EO 14110, directed NTIA to review the risks and benefits of AI models with widely available weights and develop policy recommendations to maximize those benefits while mitigating the risks.²⁵⁰ Access to model weights allows fine tuning and the removal of limitations on the model’s outputs.

At the President’s direction, NTIA sought public input through a Request for Information (RFI) public meetings, and other stakeholder engagement efforts, to ensure that the wide range of public values—from integrity of the scientific process to human rights, to innovation and competition to national security and public safety—were all considered in developing a path forward.²⁵¹ This effort is particularly important as government shaping of dual-use technology can have profound implications for the availability and values in applications and services and can hide the government’s role in determining the baked-in values and market-availability of technology making it more difficult for the public to practically and legally participate in determining technology’s shape and impact.²⁵²

The NTIA process revealed a wide range of risks and benefits associated with dual-use foundation models with widely available model weights,

247. One of Larry Lessig’s key concerns with government shaping of technology is its capacity to obscure government action from the public. LAWRENCE LESSIG, CODE AND OTHER LAWS OF CYBERSPACE 6–8 (1999).

248. Bamberger & Mulligan, *Governance-By-Design*, *supra* note 2, at 759.

249. *Id.* at 760.

250. Exec. Order. No. 14,110, *supra* note 79.

251. NAT’L TELECOMMS. & INFO. ADMIN., DUAL USE FOUNDATION ARTIFICIAL INTELLIGENCE MODELS WITH WIDELY AVAILABLE MODEL WEIGHTS 2 (July 2024), <https://www.ntia.gov/sites/default/files/publications/ntia-ai-open-model-report.pdf>.

252. Lessig, *supra* note 247 at 6–8.

including public safety, competition, research, among others.²⁵³ It broadened a conversation that had been dominated by concerns that widely available model weights would exacerbate the ability of non-experts to design, synthesize, produce, acquire, or use, chemical, biological, radiological, or nuclear (CBRN) weapons or aid individuals conducting cyberattacks. This process ensured policymakers had a more fulsome picture of the values at stake in prohibiting such models.²⁵⁴

The resulting report recommends against restrictions on the wide availability of model weights for dual-use foundation models but recommends that the U.S. government actively collect and evaluate evidence to inform future policy decisions.²⁵⁵ It evaluates a range of policy approaches, assessing their risks and benefits.²⁵⁶ And it concludes that, current evidence is not sufficient to definitively determine either that restrictions on such open weight models are warranted, or that restrictions will never be appropriate in the future.²⁵⁷ The report recommends that the government actively monitor a portfolio of risks that could arise from dual-use foundation models with widely available model weights and take steps to ensure that the government is prepared to act if heightened risks emerge.²⁵⁸

The process exemplifies the broad conversation about competing values necessary to wisely enlist design as governance. And precisely because NTIA was directed to explore the implications of technological design choices on a wide range of public rights and values, its recommendations exemplify our overarching design principle of “Design[ing] with Modesty and Restraint to Preserve Flexibility.”²⁵⁹ Because of the complex entangled values at stake, and the limited evidence of specific “marginal risks” from widely disclosed model weights in comparison to withheld model weights, NTIA recommended monitoring of the field rather than directing particular technological deployment choices—namely prohibiting the disclosure of model weights, while recommending the U.S. government undertake activities to collect, evaluate, and if appropriate act on, evidence.²⁶⁰

253. NAT’L TELECOMMS. & INFO. ADMIN., *supra* note 251, at 12–34.

254. *Id.*

255. *Id.* at 36.

256. *Id.* at 36–39.

257. *Id.* at 40–47.

258. *Id.*

259. Bamberger & Mulligan, *Governance-By-Design*, *supra* note 2, at 743.

260. NAT’L TELECOMMS. & INFO. ADMIN., *supra* note 251, at 40 (“As of the time of publication of this Report, there is not sufficient evidence on the marginal risks of dual-use foundation models with widely available model weights to conclude that restrictions on model weights are currently appropriate, nor that restrictions will never be appropriate in the future.”).

V. CONCLUSION

The Biden-Harris Administration recentered public values in AI governance. It did so through bold public statements that caught the public's imagination, acknowledged the public's experiences, and centered them—their rights and safety, and their communities—in the story of AI's future. It did so through carefully crafted bureaucratic rules that broke with traditional approaches to assessing technology, requiring agencies to assess use cases in all their messy, context-specific complexity. It did so by bringing technologists into government to support the use and governance of AI in collaboration with civil rights, civil liberties, privacy, accessibility, and other subject matter experts. It did so by creating new practices to guide technical design and use that fostered critical reflection and accountability throughout the design and deployment of systems. And it did so by requiring agencies to engage the public in these design and deployment processes and creating visibility throughout them to support public input and accountability. The Administration recognized that how we design, use, and refuse technology is a key way our nation manifests our values. And that getting technology right from the start requires technologists of many sorts to be at the table.

At a moment when rogue technologists are disregarding the rule of law, running roughshod over rights, and weaponizing technology to destroy public institutions, individuals' lives, and democracy, it is foreseeable that some will question the presence of technologists, perhaps suggesting we kick them out of the room. But that would be a mistake for many reasons. Of course, the lawyers currently shredding the government far outnumber the technologists. But more importantly, information and communication technology is a ubiquitous part of government and the institutions it regulates. Delivering robust, rights-respecting services and enforcing the laws that protect the public's rights and safety require technologists to be part of the team.