



EXECUTIVE OFFICE OF THE PRESIDENT
OFFICE OF MANAGEMENT AND BUDGET
WASHINGTON, D.C. 20503

THE DIRECTOR

March 28, 2024

M-24-10

MEMORANDUM FOR THE HEADS OF EXECUTIVE DEPARTMENTS AND AGENCIES

FROM: Shalanda D. Young 

SUBJECT: Advancing Governance, Innovation, and Risk Management for Agency Use of Artificial Intelligence

Artificial intelligence (AI) is one of the most powerful technologies of our time, and the President has been clear that we must seize the opportunities AI presents while managing its risks. Consistent with the AI in Government Act of 2020,¹ the Advancing American AI Act,² and Executive Order 14110 on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, this memorandum directs agencies to advance AI governance and innovation while managing risks from the use of AI in the Federal Government, particularly those affecting the rights and safety of the public.³

1. OVERVIEW

While AI is improving operations and service delivery across the Federal Government, agencies must effectively manage its use. As such, this memorandum establishes new agency requirements and guidance for AI governance, innovation, and risk management, including through specific minimum risk management practices for uses of AI that impact the rights and safety of the public.

Strengthening AI Governance. Managing AI risk and promoting AI innovation requires effective AI governance. As required by Executive Order 14110, each agency must designate a Chief AI Officer (CAIO) within 60 days of the date of the issuance of this memorandum. This memorandum describes the roles, responsibilities, seniority, position, and reporting structures for agency CAIOs, including expanded reporting through agency AI use case inventories. Because AI is deeply interconnected with other technical and policy areas including data, information technology (IT), security, privacy, civil rights and civil liberties, customer experience, and

¹ Pub. L. No. 116-260, div. U, title 1, § 104 (codified at 40 U.S.C. § 11301 note), <https://www.congress.gov/116/plaws/publ260/PLAW-116publ260.pdf>.

² Pub. L. No. 117-263, div. G, title LXXII, subtitle B, §§ 7224(a), 7224(d)(1)(B), and 7225 (codified at 40 U.S.C. 11301 note), <https://www.congress.gov/117/plaws/publ263/PLAW-117publ263.pdf>.

³ This memorandum accounts for public comments that OMB received following its publication of a draft version of this memorandum on November 1, 2023. OMB has separately published an explanation and response to public comments, available at <https://www.regulations.gov/document/OMB-2023-0020-0001>.

workforce management, CAIOs must work in close coordination with existing responsible officials and organizations within their agencies.

Advancing Responsible AI Innovation. With appropriate safeguards in place, AI can be a helpful tool for modernizing agency operations and improving Federal Government service to the public. To that end, agencies must increase their capacity to responsibly adopt AI, including generative AI, and take steps to enable sharing and reuse of AI models, code, and data. This memorandum requires each agency identified in the Chief Financial Officers Act (CFO Act)⁴ to develop an enterprise strategy for how they will advance the responsible use of AI. This memorandum also provides recommendations for how agencies should reduce barriers to the responsible use of AI, including barriers related to IT infrastructure, data, cybersecurity, workforce, and the particular challenges of generative AI.

Managing Risks from the Use of AI. While agencies will realize significant benefits from AI, they must also manage a range of risks from the use of AI. Agencies are subject to existing risk management requirements relevant to AI, and this memorandum does not replace or supersede these requirements. Instead, it establishes new requirements and recommendations that, both independently and collectively, address the specific risks from relying on AI to inform or carry out agency decisions and actions, particularly when such reliance impacts the rights and safety of the public. To address these risks, this memorandum requires agencies to follow minimum practices when using safety-impacting AI and rights-impacting AI, and enumerates specific categories of AI that are presumed to impact rights and safety. Finally, this memorandum also establishes a series of recommendations for managing AI risks in the context of Federal procurement.⁵

2. SCOPE

Agency adoption of AI poses many challenges, some novel and specific to AI and some well-known. While agencies must give due attention to all aspects of AI, this memorandum is more narrowly scoped to address a subset of AI risks, as well as governance and innovation issues that are directly tied to agencies' use of AI. The risks addressed in this memorandum result from any reliance on AI outputs to inform, influence, decide, or execute agency decisions or actions, which could undermine the efficacy, safety, equitableness, fairness, transparency, accountability, appropriateness, or lawfulness of such decisions or actions.⁶

⁴ 31 U.S.C. § 901(b).

⁵ Consistent with provisions of the AI in Government Act of 2020, the Advancing American AI Act, and Executive Order 14110 directing the publication of this memorandum, this memorandum sets forth multiple independent requirements and recommendations for agencies, and OMB intends that these requirements and recommendations be treated as severable. For example, the memorandum's provisions regarding the strengthening of AI governance in Section 2 are capable of operating independently, and serve an independent purpose, from the required risk management practices set forth in Section 5. Likewise, each of Section 5's individual risk management practices serves an independent purpose and can function independently from the other risk management practices. Accordingly, while this memorandum governs only agencies' own use of AI and does not create rights or obligations for the public, in the event that a court were to stay or enjoin application of a particular provision of this memorandum, or its application to a particular factual circumstance, OMB would intend that the remainder of the memorandum remain operative.

⁶ The subset of AI risks addressed in this memorandum is generally referred to in this document as "risks from the use of AI", and a full definition for this term is provided in Section 6.

This memorandum does not address issues that are present regardless of whether AI is used versus any other software, such as issues with respect to Federal information and information systems in general. In addition, this memorandum does not supersede other, more general Federal policies that apply to AI but are not focused specifically on AI, such as policies that relate to enterprise risk management, information resources management, privacy, accessibility, Federal statistical activities, IT, or cybersecurity.

Agencies must continue to comply with applicable OMB policies in other domains relevant to AI, and to coordinate compliance across the agency with all appropriate officials. All agency responsible officials retain their existing authorities and responsibilities established in other laws and policies.

a. Covered Agencies. Except as specifically noted, this memorandum applies to all agencies defined in 44 U.S.C. § 3502(1).⁷ As noted in the relevant sections, some requirements in this memorandum apply only to Chief Financial Officers Act (CFO Act) agencies as identified in 31 U.S.C. § 901(b), and other requirements do not apply to elements of the Intelligence Community, as defined in 50 U.S.C. § 3003.

b. Covered AI. This memorandum provides requirements and recommendations that, as described in more detail below, apply to new and existing AI that is developed, used, or procured by or on behalf of covered agencies. This memorandum does not, by contrast, govern:

- i. agencies' regulatory actions designed to prescribe law or policy regarding non-agency uses of AI;
- ii. agencies' evaluations of particular AI applications because the AI provider is the target or potential target of a regulatory enforcement, law enforcement, or national security action;⁸
- iii. agencies' development of metrics, methods, and standards to test and measure AI, where such metrics, methods, and standards are for use by the general public or the government as a whole, rather than to test AI for a particular agency application⁹; or
- iv. agencies' use of AI to carry out basic research or applied research, except where the purpose of such research is to develop particular AI applications within the agency.

⁷ The term "agency," as used in both the AI in Government Act of 2020 and the Advancing American AI Act, is defined as "any executive department, military department, Government corporation, Government controlled corporation, or other establishment in the executive branch of the Government (including the Executive Office of the President), or any independent regulatory agency," but does not include the Government Accountability Office; the Federal Election Commission; the governments of the District of Columbia and of the territories and possessions of the United States, and their various subdivisions; or Government-owned contractor-operated facilities, including laboratories engaged in national defense research and production activities. 44 U.S.C. § 3502(1); *see* AI in Government Act of 2020 § 102(2) (defining "agency" by reference to § 3502); Advancing American AI Act § 7223(1) (same). As a result, independent regulatory agencies as defined in 44 U.S.C. § 3502(5), which were not included in the definitions of "agency" in Executive Order 13960 and Executive Order 14110, *are* covered by this memorandum.

⁸ AI is not in scope when it is the target or potential target of such an action, but it is in scope when the AI is used to *carry out* an enforcement or national security action. For example, when evaluating an AI tool to determine whether it violates the law, the AI would not be in scope; if agencies were using that same tool to assess a different target, then the AI would be in scope.

⁹ Examples include agency actions to develop, for general use, standards or testing methodologies for evaluating or red-teaming AI capabilities.

The requirements and recommendations of this memorandum apply to system functionality that implements or is reliant on AI, rather than to the entirety of an information system that incorporates AI. As noted in the relevant sections, some requirements in this memorandum apply only in specific circumstances in which agencies use AI, such as when the AI may impact rights or safety.

c. Applicability to National Security Systems. This memorandum does not cover AI when it is being used as a component of a National Security System.¹⁰

3. STRENGTHENING ARTIFICIAL INTELLIGENCE GOVERNANCE

The head of each covered agency is responsible for pursuing AI innovation and ensuring that their agency complies with AI requirements in relevant law and policy, including the requirement that risks from the agency's use of AI are adequately managed. Doing so requires a strong governance structure and agencies are encouraged to strategically draw upon their policy, programmatic, research and evaluation, and regulatory functions to support the implementation of this memorandum's requirements and recommendations. The head of each covered agency must also consider the financial, human, information, and infrastructure resources necessary for implementation, prioritizing current resources or requesting additional resources via the budget process, as needed to support the responsibilities identified in this memorandum.

To improve accountability for AI issues, agencies must designate a Chief AI Officer, consistent with Section 10.1(b) of Executive Order 14110. CAIOs bear primary responsibility on behalf of the head of their agency for implementing this memorandum and coordinating implementation with other agencies. This section defines CAIOs' roles, responsibilities, seniority, position, and reporting structure.

a. Actions

- i. **Designating Chief AI Officers.** Within 60 days of the issuance of this memorandum, the head of each agency must designate a CAIO. To ensure the CAIO can fulfill the responsibilities laid out in this memorandum, agencies that have already designated a CAIO must evaluate whether they need to provide that individual with additional authority or appoint a new CAIO. Agencies must identify these officers to OMB through OMB's Integrated Data Collection process or an OMB-designated successor process. When the designated individual changes or the position is vacant, agencies must notify OMB within 30 days.
- ii. **Convening Agency AI Governance Bodies.** Within 60 days of the issuance of this memorandum, each CFO Act agency must convene its relevant senior officials to

¹⁰ AI innovation and risk for National Security Systems must be managed appropriately, but these systems are governed through other policy. For example, Section 4.8 of Executive Order 14110 directs the development of a National Security Memorandum to govern the use of AI as a component of a National Security System, and agencies also have existing guidelines in place, such as the Department of Defense's (DoD) *Responsible Artificial Intelligence Strategy and Implementation Pathway* and the Office of the Director of National Intelligence's *Principles of Artificial Intelligence Ethics for the Intelligence Community*, as well as policies governing specific high-risk national security applications of AI, such as DoD Directive 3000.09, *Autonomy in Weapon Systems*.

coordinate and govern issues tied to the use of AI within the Federal Government, consistent with Section 10.1(b) of Executive Order 14110 and the detailed guidance in Section 3(c) of this memorandum.

- iii. **Compliance Plans.** Consistent with Section 104(c) and (d) of the AI in Government Act of 2020, within 180 days of the issuance of this memorandum or any update to this memorandum, and every two years thereafter until 2036, each agency must submit to OMB and post publicly on the agency’s website either a plan to achieve consistency with this memorandum, or a written determination that the agency does not use and does not anticipate using covered AI. Agencies must also include plans to update any existing internal AI principles and guidelines to ensure consistency with this memorandum.¹¹ OMB will provide templates for these compliance plans.
- iv. **AI Use Case Inventories.** Each agency (except for the Department of Defense and the Intelligence Community) must individually inventory each of its AI use cases at least annually, submit the inventory to OMB, and post a public version on the agency’s website. OMB will issue detailed instructions for the inventory and its scope through its Integrated Data Collection process or an OMB-designated successor process. Beginning with the use case inventory for 2024, agencies will be required, as applicable, to identify which use cases are safety-impacting and rights-impacting AI and report additional detail on the risks—including risks of inequitable outcomes—that such uses pose and how agencies are managing those risks.
- v. **Reporting on AI Use Cases Not Subject to Inventory.** Some AI use cases are not required to be individually inventoried, such as those in the Department of Defense or those whose sharing would be inconsistent with applicable law and governmentwide policy. On an annual basis, agencies must still report and release aggregate metrics about such use cases that are otherwise within the scope of this memorandum, the number of such cases that impact rights and safety, and their compliance with the practices of Section 5(c) of this memorandum. OMB will issue detailed instructions for this reporting through its Integrated Data Collection process or an OMB-designated successor process.

b. Roles, Responsibilities, Seniority, Position, and Reporting Structure of Chief Artificial Intelligence Officers

Consistent with Section 10.1(b)(ii) of Executive Order 14110, this memorandum defines CAIOs’ roles, responsibilities, seniority, position, and reporting structures as follows:

- i. **Roles.** CAIOs must have the necessary skills, knowledge, training, and expertise to perform the responsibilities described in this section. At CFO Act agencies, a primary role of the CAIO must be coordination, innovation, and risk management for their agency’s use of AI specifically, as opposed to data or IT issues in general. Agencies may choose to designate an existing official, such as a Chief Information Officer (CIO), Chief Data Officer (CDO), Chief Technology Officer, or similar official with relevant or

¹¹ Given the importance of context-specific guidance on AI, agencies are encouraged to continue implementing their agency’s AI principles and guidelines, so long as they do not conflict with this memorandum.

complementary authorities and responsibilities, provided they have significant expertise in AI and meet the other requirements in this section.

- ii. **Responsibilities.** Executive Order 14110 tasks CAIOs with primary responsibility in their agencies, in coordination with other responsible officials, for coordinating their agency's use of AI, promoting AI innovation, managing risks from the use of AI, and carrying out the agency responsibilities defined in Section 8(c) of Executive Order 13960¹² and Section 4(b) of Executive Order 14091.¹³ In addition, CAIOs, in coordination with other responsible officials and appropriate stakeholders, are responsible for:

Coordinating Agency Use of AI

- A. serving as the senior advisor for AI to the head of the agency and other senior agency leadership and within their agency's senior decision-making forums;
- B. instituting the requisite governance and oversight processes to achieve compliance with this memorandum and enable responsible use of AI in the agency, in coordination with relevant agency officials;
- C. maintaining awareness of agency AI activities, including through the creation and maintenance of the annual AI use case inventory;
- D. developing a plan for compliance with this memorandum, as detailed in Section 3(a)(iii) of this memorandum, and an agency AI strategy, as detailed in Section 4(a) of this memorandum;
- E. working with and advising the agency CFO on the resourcing requirements necessary to implement this memorandum and providing recommendations on priority investment areas to build upon existing enterprise capacity;
- F. advising the Chief Human Capital Officer (CHCO) and where applicable, the Chief Learning Officer, on improving workforce capacity and securing and maintaining the skillsets necessary for using AI to further the agency's mission and adequately manage its risks;
- G. sharing relevant information with agency officials involved in the agency's major AI policymaking initiatives;
- H. supporting agency involvement with appropriate interagency coordination bodies related to their agency's AI activities, including representing the agency on the council described in Section 10.1(a) of Executive Order 14110;
- I. supporting and coordinating their agency's involvement in AI standards-setting bodies, as appropriate, and encouraging agency adoption of voluntary consensus standards for AI, as appropriate and consistent with OMB Circular No. A-119, if applicable;¹⁴

¹² Executive Order 13960, *Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government*, <https://www.govinfo.gov/content/pkg/FR-2020-12-08/pdf/2020-27065.pdf>.

¹³ Executive Order 14091, *Further Advancing Racial Equity and Support for Underserved Communities Through the Federal Government*, <https://www.govinfo.gov/content/pkg/FR-2023-02-22/pdf/2023-03779.pdf>.

¹⁴ OMB Circular A-119, *Federal Participation in the Development and Use of Voluntary Consensus Standards and in Conformity Assessment Activities* (Feb. 10, 1998), <https://www.whitehouse.gov/wp-content/uploads/2017/11/Circular-119-1.pdf>.

- J. promoting equity and inclusion within the agency's AI governance structures and incorporating diverse perspectives into the decision-making process;

Promoting AI Innovation

- K. working with their agency to identify and prioritize appropriate uses of AI that will advance both their agency's mission and equitable outcomes;
- L. identifying and removing barriers to the responsible use of AI in the agency, including through the advancement of AI-enabling enterprise infrastructure, data access and governance, workforce development measures, policy, and other resources for AI innovation;
- M. working with their agency's CIO, CDO, and other relevant officials to ensure that custom-developed AI code and the data used to develop and test AI are appropriately inventoried, shared, and released in agency code and data repositories in accordance with Section 4(d) of this memorandum;
- N. advocating within their agency and to the public on the opportunities and benefits of AI to the agency's mission;

Managing Risks from the Use of AI

- O. managing an agency program that supports the enterprise in identifying and managing risks from the use of AI, especially for safety-impacting and rights-impacting AI;
- P. working with relevant senior agency officials to establish or update processes to measure, monitor, and evaluate the ongoing performance and effectiveness of the agency's AI applications and whether the AI is advancing the agency's mission and meeting performance objectives;
- Q. overseeing agency compliance with requirements to manage risks from the use of AI, including those established in this memorandum and in relevant law and policy;
- R. conducting risk assessments, as necessary, of the agency's AI applications to ensure compliance with this memorandum;
- S. working with relevant agency officials to develop supplementary AI risk management guidance particular to the agency's mission, including working in coordination with officials responsible for privacy and civil rights and civil liberties on identifying safety-impacting and rights-impacting AI within the agency;
- T. waiving individual applications of AI from elements of Section 5 of this memorandum through the processes detailed in that section; and
- U. in partnership with relevant agency officials (e.g., authorizing, procurement, legal, data governance, human capital, and oversight officials), establishing controls to ensure that their agency does not use AI that is not in compliance with this memorandum, including by assisting these relevant agency officials in evaluating Authorizations to Operate based on risks from the use of AI.

- iii. **Seniority.** For CFO Act agencies, the CAIO must be a position at the Senior Executive Service, Scientific and Professional, or Senior Leader level, or equivalent. In other agencies, the CAIO must be at least a GS-15 or equivalent.
- iv. **Position and Reporting Structure.** CAIOs must have the necessary authority to perform the responsibilities in this section and must be positioned highly enough to engage regularly with other agency leadership, to include the Deputy Secretary or equivalent. Further, CAIOs must coordinate with other responsible officials at their agency to ensure that the agency's use of AI complies with and is appropriate in light of applicable law and governmentwide guidance.

c. Internal Agency AI Coordination

Agencies must ensure that AI issues receive adequate attention from the agency's senior leadership. Consistent with Section 10.1(b) of Executive Order 14110, agencies must take appropriate steps, such as through the convening of an AI governance body, to coordinate internally among officials responsible for aspects of AI adoption and risk management. Likewise, the CAIO must be involved, at appropriate times, in broader agency-wide risk management bodies and processes,¹⁵ including in the development of the agency risk management strategy.¹⁶ The agency's AI coordination mechanisms should be aligned to the needs of the agency based on, for example, the degree to which the agency currently uses AI, the degree to which AI could improve the agency's mission, and the risks posed by the agency's current and potential uses of AI.

Each CFO Act agency is required to establish an AI Governance Board to convene relevant senior officials to govern the agency's use of AI, including to remove barriers to the use of AI and to manage its associated risks. Those agencies are permitted to rely on existing governance bodies¹⁷ to fulfill this requirement as long as they currently satisfy or are made to satisfy both of the following:

- i. Agency AI Governance Boards must be chaired by the Deputy Secretary of the agency or equivalent and vice-chaired by the agency CAIO, and these roles should not be assigned to other officials. The full Board, including the Deputy Secretary, must convene on at least a semi-annual basis. Working through this Board, CAIOs will support their respective Deputy Secretaries in coordinating AI activities across the agency and implementing relevant sections of Executive Order 14110.

¹⁵ See, e.g., OMB Circular No. A-123, *Management's Responsibility for Enterprise Risk Management and Internal Control* (July 15, 2016), https://www.whitehouse.gov/wp-content/uploads/legacy_drupal_files/omb/memoranda/2016/m-16-17.pdf.

¹⁶ See OMB Circular No. A-130, *Managing Information as a Strategic Resource*, Appx. I, sec. 5(b) (July 28, 2016), https://www.whitehouse.gov/wp-content/uploads/legacy_drupal_files/omb/circulars/A130/a130revised.pdf.

¹⁷ An example of a qualifying body includes agency Data Governance Bodies, established by OMB Memorandum M-19-23, *Phase I Implementation of the Foundations for Evidence-Based Policymaking Act of 2018: Learning Agendas, Personnel, and Planning Guidance*, <https://www.whitehouse.gov/wp-content/uploads/2019/07/m-19-23.pdf>.

- ii. Agency AI Governance Boards must include appropriate representation from senior agency officials responsible for key enablers of AI adoption and risk management, including at least IT, cybersecurity, data, privacy, civil rights and civil liberties, equity, statistics, human capital, procurement, budget, legal, agency management, customer experience, program evaluation, and officials responsible for implementing AI within an agency's program office(s). Agencies should also consider including representation from their respective Office of the Inspector General.

Agencies are encouraged to have their AI Governance Boards consult external experts as appropriate and consistent with applicable law. Experts' individual viewpoints can help broaden the perspective of an existing governance board and inject additional technical, ethics, civil rights and civil liberties, or sector-specific expertise, as well as methods for engaging the workforce.

4. ADVANCING RESPONSIBLE ARTIFICIAL INTELLIGENCE INNOVATION

If implemented responsibly, AI can improve operations and deliver efficiencies across the Federal Government. Agencies must improve their ability to use AI in ways that benefit the public and increase mission effectiveness, while recognizing the limitations and risks of AI and when it is not suited for a given task. In particular, agencies are encouraged to prioritize AI development and adoption for the public good and where the technology can be helpful in understanding and tackling large societal challenges, such as using AI to improve the accessibility of government services, reduce food insecurity, address the climate crisis, improve public health, advance equitable outcomes, protect democracy and human rights, and grow economic competitiveness in a way that benefits people across the United States.

To achieve this, agencies should build upon existing internal enterprise capacity to support responsible AI innovation, take actions to strengthen their AI and AI-enabling talent,¹⁸ and improve their ability to develop and procure AI. Agencies should both explore joint efforts to scale these opportunities as well as take steps to responsibly share their AI resources across the Federal Government and with the public.

a. AI Strategies

Within 365 days of the issuance of this memorandum, each CFO Act agency must develop and release publicly on the agency's website a strategy for identifying and removing barriers to the responsible use of AI and achieving enterprise-wide improvements in AI maturity, including:

- i. the agency's current and planned uses of AI that are most impactful to an agency's mission or service delivery;¹⁹

¹⁸ Agencies should also ensure that they consider and satisfy applicable collective bargaining obligations regarding their implementation of AI.

¹⁹ Consistent with Sections 7225(d) and 7228 of the Advancing American AI Act, this requirement applies to CFO Act agencies except for the Department of Defense, and does not apply to elements of the Intelligence Community,

- ii. a current assessment of the agency's AI maturity and the agency's AI maturity goals;
- iii. the agency's plans to effectively govern its use of AI, including through its Chief AI Officer, AI Governance Boards, and improvements to its AI use case inventory;
- iv. a plan for developing sufficient enterprise capacity for AI innovation, including mature AI-enabling infrastructure for the data, computing, development, testing, cybersecurity compliance, deployment, and continuous-monitoring infrastructure necessary to build, test, and maintain AI;
- v. a plan for providing sufficient AI tools and capacity to support the agency's research and development (R&D) work consistent with the R&D priorities developed by OMB and the Office of Science and Technology Policy, the National AI R&D Strategic Plan, and agency-specific R&D plans;
- vi. a plan for establishing operational and governance processes as well as developing the necessary infrastructure to manage risks from the use of AI;
- vii. a current assessment of the agency's AI and AI-enabling workforce capacity and projected AI and AI-enabling workforce needs, as well as a plan to recruit, hire, train, retain, and empower AI practitioners and achieve AI literacy for non-practitioners involved in AI to meet those needs;
- viii. the agency's plan to encourage diverse perspectives throughout the AI development or procurement lifecycle, including how to determine whether a particular use of AI is meeting the agency's equity goals and civil rights commitments; and
- ix. specific, prioritized areas and planning for future AI investment, leveraging the annual budget process as appropriate.

b. Removing Barriers to the Responsible Use of AI

Embracing innovation requires removing unnecessary and unhelpful barriers to the use of AI while retaining and strengthening the guardrails that ensure its responsible use. Agencies should create internal environments where those developing and deploying AI have sufficient flexibility and where limited AI resources and expertise are not diverted away from AI innovation and risk management. Agencies should take steps to remove barriers to responsible use of AI, paying special attention to the following recommendations:

- i. **IT Infrastructure.** Agencies should ensure that their AI projects have access to adequate IT infrastructure, including high-performance computing infrastructure specialized for AI training and inference, where necessary. Agencies should also ensure adequate access for AI developers to the software tools, open-source libraries, and deployment and

as defined in 50 U.S.C. § 3003(4). Information that would be protected from release if requested under 5 U.S.C. § 552 need not be included in the strategy.

monitoring capabilities necessary to rapidly develop, test, and maintain AI applications.

- ii. **Data.** Agencies should develop adequate infrastructure and capacity to sufficiently share, curate, and govern agency data for use in training, testing, and operating AI. This includes an agency's capacity to maximize appropriate access to and sharing of both internally held data and agency data managed by third parties. Agencies should also explore the possible utility of and legal authorities supporting the use of publicly available information, and encourage its use where appropriate and consistent with the data practices outlined in this memorandum. Any data used to help develop, test, or maintain AI applications, regardless of source, should be assessed for quality, representativeness, and bias. These activities should be supported by resources to enable sound data governance and management practices, particularly as they relate to data collection, curation, labeling, and stewardship.
- iii. **Cybersecurity.** Agencies should update, as necessary, processes for information system authorization and continuous monitoring to better address the needs of AI applications, including to advance the use of continuous authorizations for AI. Consistent with Section 10.1(f) of Executive Order 14110, agency authorizing officials are encouraged to prioritize review of generative AI and other critical and emerging technologies in Authorizations to Operate and any other applicable release or oversight processes.
- iv. **Generative AI.** In addition to following the guidance provided in Section 10.1(f) of Executive Order 14110, agencies should assess potential beneficial uses of generative AI in their missions and establish adequate safeguards and oversight mechanisms that allow generative AI to be used in the agency without posing undue risk.

c. AI Talent

Consistent with Section 10.2 of Executive Order 14110, agencies are strongly encouraged to prioritize recruiting, hiring, developing, and retaining talent in AI and AI-enabling roles to increase enterprise capacity for responsible AI innovation. Agencies should:

- i. follow the hiring practices described in the forthcoming AI and Tech Hiring Playbook created by the Office of Personnel Management (OPM), including encouraging applications from individuals with diverse perspectives, making best use of available hiring and retention authorities and using descriptive job titles and skills-based assessments;
- ii. designate an AI Talent Lead who, for at least the duration of the AI Talent Task Force, will be accountable for reporting to agency leadership, tracking AI hiring across the agency, and providing data to OPM and OMB on hiring needs and progress. The AI Talent Task Force, established in Section 10.2(b) of EO 14110, will provide AI Talent Leads with engagement opportunities to enhance their AI hiring practices and to drive impact through collaboration across agencies, including sharing position descriptions, coordinating marketing and outreach, shared hiring actions, and, if appropriate, sharing

applicant information across agencies; and

- iii. in consultation with Federal employees and their union representatives, where applicable, provide resources and training to develop AI talent internally and increase AI training offerings for Federal employees, including opportunities that provide Federal employees pathways to AI occupations and that assist employees affected by the application of AI to their work.

d. AI Sharing and Collaboration

Openness, sharing, and reuse of AI significantly enhance both innovation and transparency, and must also be done responsibly to avoid undermining the rights, safety, and security of the public. Agencies must share their AI code, models, and data, and do so in a manner that facilitates re-use and collaboration Government-wide and with the public, subject to applicable law, governmentwide guidance, and the following considerations:

- i. **Sharing and Releasing AI Code and Models.** Agencies must proactively share their custom-developed code²⁰—including models and model weights—for AI applications in active use and must release and maintain that code as open source software on a public repository,²¹ unless:
 - A. the sharing of the code is restricted by law or regulation, including patent or intellectual property law, the Export Asset Regulations, the International Traffic in Arms Regulations, and Federal laws and regulations governing classified information;
 - B. the sharing of the code would create an identifiable risk to national security, confidentiality of Government information, individual privacy, or the rights or safety of the public;
 - C. the agency is prevented by a contractual obligation from doing so; or
 - D. the sharing of the code would create an identifiable risk to agency mission, programs, or operations, or to the stability, security, or integrity of an agency’s systems or personnel.

Agencies should prioritize sharing custom-developed code, such as commonly used packages or functions, that has the greatest potential for re-use by other agencies or the public.

- ii. **Sharing and Releasing AI Data Assets.** Data used to develop and test AI is likely to constitute a “data asset” for the purposes of implementing the Open, Public, Electronic

²⁰ A full definition for “custom-developed code” is provided in Section 6.

²¹ For guidance and best practices related to sharing code and releasing it as open source, agencies should consult OMB Memorandum M-16-21, *Federal Source Code Policy: Achieving Efficiency, Transparency, and Innovation through Reusable and Open Source Software* (Aug. 8, 2016), https://www.whitehouse.gov/wp-content/uploads/legacy_drupal_files/omb/memoranda/2016/m_16_21.pdf. Agencies are additionally encouraged to draw upon existing collaboration methods to facilitate the sharing and release of code and models, including the council described in Section 10.1(a) of Executive Order 14110, the General Services Administration’s AI Community of Practice, and <https://www.code.gov>, as well as other publicly available code repositories.

and Necessary (OPEN) Government Data Act,²² and agencies must, if required by that Act and pursuant to safety and security considerations in Section 4.7 of Executive Order 14110, release such data assets publicly as open government data assets.²³ When sharing AI data assets, agencies should promote data interoperability, including by coordinating internally and with other relevant agencies on interoperability criteria and using standardized data formats where feasible and appropriate.

- iii. **Partial Sharing and Release.** Where some portion of an AI project’s code, models, or data cannot be shared or released publicly pursuant to subsections (i) and (ii) of this section, the rest should still be shared or released where practicable, such as by releasing the data used to evaluate a model even if the model itself cannot be safely released, or by sharing a model within the Federal Government even if the model cannot be publicly released. Where code, models, or data cannot be released without restrictions on who can access it, agencies should also, where practicable, share them through Federally controlled infrastructure that allows controlled access by entities outside the Federal Government, such as via the National AI Research Resource.
- iv. **Procuring AI for Sharing and Release.** When procuring custom-developed code for AI, data to train and test AI, and enrichments to existing data (such as labeling services), agencies are encouraged to do so in a manner that allows for the sharing and public release of the relevant code, models, and data.
- v. **Unintended Disclosure of Data from AI Models.** When agencies are deciding whether to share and release AI models and model weights, they should assess the risk that the models can be induced to reveal sensitive details of the data used to develop them. Agencies’ assessment of risk should include a model-specific risk analysis.²⁴

e. Harmonization of Artificial Intelligence Requirements

Interpreting and implementing AI management requirements in a consistent manner across Federal agencies will create efficiencies as well as opportunities for sharing resources and best practices. To assist in this effort and consistent with Section 10.1(a) of Executive Order 14110, OMB, in collaboration with the Office of Science and Technology Policy, will coordinate the development and use of AI in agencies’ programs and operations—including the implementation of this memorandum—across Federal agencies through an interagency council. This will include at a minimum:

- i. promoting shared templates and formats;
- ii. sharing best practices and lessons learned, including for achieving meaningful participation from affected communities and the public in AI development and

²² Title II of the Foundations for Evidence-Based Policymaking Act of 2018, P.L. 115-435.

²³ Where such data is already publicly available, agencies are not required to duplicate it, but should maintain and share the provenance of such data and how others can access it.

²⁴ The risks of unintended disclosure differ by model, and agencies should also not assume that an AI model poses the same privacy and confidentiality risks as the data used to develop it.

procurement, updating organizational processes to better accommodate AI, removing barriers to responsible AI innovation, responding to AI incidents that may have resulted in harm to an individual, and building a diverse AI workforce to meet the agency's needs;

- iii. sharing technical resources for implementation of this memorandum's risk management practices, such as for testing, continuous monitoring, and evaluation; and
- iv. highlighting exemplary uses of AI for agency adoption, particularly uses which help address large societal challenges.

5. MANAGING RISKS FROM THE USE OF ARTIFICIAL INTELLIGENCE

Agencies have a range of policies, procedures, and officials in place to manage risks related to agency information and systems. To better address risks from the use of AI, and particularly risks to the rights and safety of the public, all agencies are required to implement minimum practices, detailed below, to manage risks from safety-impacting AI and rights-impacting AI.²⁵ However, Section 5(a) through (c) of this memorandum do not apply to elements of the Intelligence Community.²⁶

a. Actions

- i. **Implementation of Risk Management Practices and Termination of Non-Compliant AI.** By December 1, 2024, agencies must implement the minimum practices in Section 5(c) of this memorandum for safety-impacting and rights-impacting AI, or else stop using any AI in their operations that is not compliant with the minimum practices, consistent with the details and caveats in that section.
- ii. **Certification and Publication of Determinations and Waivers.** By December 1, 2024, and annually thereafter, each agency must certify the ongoing validity of the determinations made under subsection (b) and the waivers granted under subsection (c) of this section. To the extent consistent with law and governmentwide policy, the agency must publicly release a summary detailing each individual determination and waiver, as well its justification. Alternatively, if an agency has no active determinations or waivers, it must publicly indicate that fact and report it to OMB. OMB will issue detailed instructions for these summaries through its Integrated Data Collection process or an OMB-designated successor process.

²⁵ Agencies are not required to incorporate these practices into criteria for granting federal financial assistance (FFA). However, they are encouraged, consistent with applicable law, to consider the minimum practices when choosing such criteria.

²⁶ Although elements of the Intelligence Community are not required to implement these practices, they are encouraged to do so.

b. Determining Which Artificial Intelligence Is Presumed to Be Safety-Impacting or Rights-Impacting

All AI that matches the definitions of “safety-impacting AI” or “rights-impacting AI” as defined in Section 6 must follow the minimum practices in Section 5(c) by the applicable deadline. Agencies must review each current or planned use of AI to assess whether it matches the definition of safety-impacting AI or rights-impacting AI. When conducting such an assessment, as reflected by the definitions of safety-impacting AI and rights-impacting AI in Section 6 of this memorandum, agencies must look to whether the particular AI output serves as a principal basis for a decision or action.

Additionally, AI used for one of the purposes identified in Appendix I is automatically *presumed* to be safety-impacting or rights-impacting. However, the agency CAIO, in coordination with other relevant officials, may determine (or revisit a prior determination) that a particular AI application or component²⁷ subject to this presumption does not match the definitions of “safety-impacting AI” or “rights-impacting AI” and is therefore not subject to the minimum practices. The agency CAIO may make or revisit such a determination only with a documented context-specific and system-specific risk assessment and may revisit a prior determination at any time. This responsibility shall not be delegated to other officials. In addition to the certification and publication requirements in Section 5(a)(ii) of this memorandum, CAIOs must centrally track these determinations, reassess them if there are significant changes to the conditions or context in which the AI is used, and report to OMB within 30 days of making or changing a determination, detailing the scope, justification, and supporting evidence.

c. Minimum Practices for Safety-Impacting and Rights-Impacting Artificial Intelligence

Except as prevented by applicable law and governmentwide guidance, agencies must apply the minimum risk management practices in this section to safety-impacting and rights-impacting AI by December 1, 2024, or else stop using the AI until they achieve compliance. Prior to December 1, 2024, agency CAIOs should work with their agencies’ relevant officials to bring potentially non-compliant AI into conformity, which may include requests that third-party vendors voluntarily take appropriate action (e.g., via updated documentation or testing measures). To ensure compliance with this requirement, relevant agency officials must use existing mechanisms wherever possible, (for example, the Authorization to Operate process).²⁸ An agency may also request an extension or grant a waiver to this requirement through its CAIO using the processes detailed below.

²⁷ CAIOs may also make these determinations across groups of AI applications or components that are closely related by design or deployment context, provided that: (1) those systems have undergone a risk assessment that adequately considers the risks from each individual system or from all possible systems in the group; and (2) the systems are substantially identical in their risk profiles.

²⁸ While agencies must use existing authorization and oversight processes to enforce these practices, the practices are most effective when applied early in the research, design, and development of AI systems, and agencies should plan for and adopt the practices throughout the relevant AI systems’ lifecycles and as early as possible, as appropriate.

Agencies must document their implementation of these practices and be prepared to report them to OMB, either as a component of the annual AI use case inventory, periodic accountability reviews, or upon request as determined by OMB.

The practices in this section represent an initial baseline for managing risk from the use of AI. Agencies must identify additional context-specific risks that are associated with their use of AI and address them as appropriate. Such risk considerations may include impacts to safety, security, civil rights, civil liberties, privacy, democratic values, human rights, equal opportunities, worker well-being, access to critical resources and services, agency trust and credibility, and market competition. To address these potential risk management gaps, agencies are encouraged to promote and to incorporate, as appropriate, additional best practices for AI risk management, such as from the National Institute of Standards and Technology (NIST) AI Risk Management Framework,²⁹ the Blueprint for an AI Bill of Rights,³⁰ relevant international standards,³¹ and the workforce principles and best practices for employers established pursuant to Section 6(b)(i) of Executive Order 14110. Agencies are also encouraged to continue developing their own agency-specific practices, as appropriate and consistent with this memorandum and the principles in Executive Order 13960, Executive Order 14091, and Executive Order 14110.

The practices in this section also do not supersede, modify, or direct an interpretation of existing requirements mandated by law or governmentwide policy, and responsible agency officials must coordinate to ensure that the adoption of these practices does not conflict with other applicable law or governmentwide guidance.

- i. **Exclusions from Minimum Practices.** Agencies are not required to follow the minimum practices outlined in this section when using AI *solely* to:
 - A. evaluate a potential vendor, commercial capability, or freely available AI capability that is not otherwise used in agency operations, exclusively for the purpose of making a procurement or acquisition decision; or
 - B. achieve its conformity with the requirements of this section, such as using an AI application in controlled testing conditions to carry out the minimum testing requirements below.³²
- ii. **Extensions for Minimum Practices.** Agencies may request from OMB an extension of up to one year, for a particular use of AI that cannot feasibly meet the minimum requirements in this section by that date. OMB will not grant renewals beyond the initial one-year extension. Any extension requests shall be submitted prior to October 15, 2024. The request must be accompanied by a detailed justification for why the agency cannot achieve compliance for the use of AI in question and what practices the agency has in

²⁹ *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST Publication AI 100-1, <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>.

³⁰ *Blueprint for an AI Bill of Rights*, White House Office of Science and Technology Policy, <https://www.whitehouse.gov/wp-content/uploads/2022/10/Blueprint-for-an-AI-Bill-of-Rights.pdf>.

³¹ For example, ISO/IEC 23894:2023 Information technology — Artificial intelligence — Guidance on risk management, <https://www.iso.org/standard/77304.html>.

³² This exclusion must not be applied to any use of AI in real-world conditions, except as specifically allowed by this section.

place to mitigate the risks from noncompliance, as well as a plan for how the agency will come to implement the full set of required minimum practices from this section. OMB will issue detailed instructions for extension requests through its Integrated Data Collection process or an OMB-designated successor process.

- iii. **Waivers from Minimum Practices.** In coordination with other relevant officials, an agency CAIO may waive one or more of the requirements in this section for a specific covered AI application or component³³ after making a written determination, based upon a system-specific and context-specific risk assessment, that fulfilling the requirement would increase risks to safety or rights overall or would create an unacceptable impediment to critical agency operations. An agency CAIO may also revoke a previously issued waiver at any time. This responsibility shall not be delegated to other officials. In addition to the certification and publication requirements in Section 5(a)(ii) of this memorandum, CAIOs must centrally track waivers, reassess them if there are significant changes to the conditions or context in which the AI is used, and report to OMB within 30 days of granting or revoking any waiver, detailing the scope, justification, and supporting evidence.
- iv. **Minimum Practices for Either Safety-Impacting or Rights-Impacting AI.** No later than December 1, 2024, agencies must follow these practices *before* using new or existing covered safety-impacting or rights-impacting AI:
 - A. **Complete an AI impact assessment.** Agencies should update their impact assessments periodically and leverage them throughout the AI's lifecycle. In their impact assessments, agencies must document at least the following:
 - 1. *The intended purpose for the AI and its expected benefit*, supported by specific metrics or qualitative analysis. Metrics should be quantifiable measures of positive outcomes for the agency's mission—for example to reduce costs, wait time for customers, or risk to human life—that can be measured using performance measurement or program evaluation methods after the AI is deployed to demonstrate the value of using AI.³⁴ Where quantification is not feasible, qualitative analysis should demonstrate an expected positive outcome, such as for improvements to customer experience, and it should demonstrate that AI is better suited to accomplish the relevant task as compared to alternative strategies.
 - 2. *The potential risks of using AI*, as well as what, if any, additional mitigation measures, beyond these minimum practices, the agency will take to help

³³ CAIOs may also grant waivers applicable to groups of AI applications or components that are closely related by design or deployment context, provided that: (1) those systems have undergone a risk assessment that adequately considers the risks from each individual system or from all possible systems in the group; and (2) the systems are substantially identical in their risk profiles.

³⁴ For supervised and semi-supervised AI, agencies should use a target variable which can be reliably measured and adequately represents the desired real-world outcomes.

reduce these risks. Agencies should document the stakeholders³⁵ who will be most impacted by the use of the system and assess the possible failure modes of the AI and of the broader system, both in isolation and as a result of human users and other likely variables outside the scope of the system itself. Agencies should be especially attentive to the potential risks to underserved communities. The expected benefits of the AI functionality should be considered against its potential risks, and if the benefits do not meaningfully outweigh the risks, agencies should not use the AI.

3. *The quality and appropriateness of the relevant data.* Agencies must assess the quality of the data used in the AI's design, development, training, testing, and operation and its fitness to the AI's intended purpose. In conducting assessments, if the agency cannot obtain such data after a reasonable effort to do so, it must obtain sufficient descriptive information from the vendor (e.g., AI or data provider) to satisfy the reporting requirements in this paragraph. At a minimum, agencies must document:
 - a. the data collection and preparation process, which must also include the provenance of any data used to train, fine-tune, or operate the AI;
 - b. the quality³⁶ and representativeness³⁷ of the data for its intended purpose;
 - c. how the data is relevant to the task being automated and may reasonably be expected to be useful for the AI's development, testing, and operation;
 - d. whether the data contains sufficient breadth to address the range of real-world inputs the AI might encounter and how data gaps and shortcomings have been addressed either by the agency or vendor; and
 - e. if the data is maintained by the Federal Government, whether that data is publicly disclosable as an open government data asset, in accordance with applicable law and policy.³⁸

- B. **Test the AI for performance in a real-world context.** Agencies must conduct adequate testing to ensure the AI, as well as components that rely on it, will work in its intended real-world context. Such testing should follow domain-specific best practices, when available, and should take into account both the specific technology used and feedback from human operators, reviewers, employees, and

³⁵ Stakeholders will vary depending on how AI is being used. For example, if an agency is using AI to control a water treatment process, stakeholders may include (1) local residents; (2) state, local, tribal, and territorial government representatives; and (3) environmental experts.

³⁶ Consistent with OMB Memorandum M-19-15, *Improving Implementation of the Information Quality Act*, <https://www.whitehouse.gov/wp-content/uploads/2019/04/M-19-15.pdf>, if applicable. Agencies should also consider the National Science and Technology Council's report *Protecting the Integrity of Government Science*, https://www.whitehouse.gov/wp-content/uploads/2022/01/01-22-Protecting_the_Integrity_of_Government_Science.pdf.

³⁷ Agencies should assess whether the data used can produce or amplify inequitable outcomes as a result of poor data representativeness or harmful bias. Such outcomes can result from historical discrimination, such as the perpetuation of harmful gender-based and racial stereotypes in society.

³⁸ See 44 U.S.C. § 3502(20).

customers who use the service or are impacted by the system's outcomes. Testing conditions should mirror as closely as possible the conditions in which the AI will be deployed. Through test results, agencies should demonstrate that the AI will achieve its expected benefits and that associated risks will be sufficiently mitigated, or else the agency should not use the AI. In conducting such testing, if an agency does not have access to the underlying source code, models, or data, the agency must use alternative test methodologies, such as querying the AI service and observing the outputs or providing evaluation data to the vendor and obtaining results. Agencies are also encouraged to leverage pilots and limited releases, with strong monitoring, evaluation, and safeguards in place, to carry out the final stages of testing before a wider release.

- C. **Independently evaluate the AI.** Agencies, through the CAIO, an agency AI oversight board, or other appropriate agency office with existing test and evaluation responsibilities, must review relevant AI documentation to ensure that the system works appropriately and as intended, and that its expected benefits outweigh its potential risks. At a minimum, this documentation must include the completed impact assessment and results from testing AI performance in a real-world context referenced in paragraphs (A) and (B) of this subsection. Agencies must incorporate this independent evaluation into an applicable release or oversight process, such as the Authorization to Operate process. The independent reviewing authority must not have been directly involved in the system's development.

No later than December 1, 2024 and on an ongoing basis *while* using new or existing covered safety-impacting or rights-impacting AI, agencies must ensure these practices are followed for the AI:

- D. **Conduct ongoing monitoring.** In addition to pre-deployment testing, agencies must institute ongoing procedures to monitor degradation of the AI's functionality and to detect changes in the AI's impact on rights and safety. Agencies should also scale up the use of new or updated AI features incrementally where possible to provide adequate time to monitor for adverse performance or outcomes. Agencies should monitor and defend the AI from AI-specific exploits,³⁹ particularly those that would adversely impact rights and safety.
- E. **Regularly evaluate risks from the use of AI.** The monitoring process in paragraph (D) must include periodic human reviews to determine whether the deployment context, risks, benefits, and agency needs have evolved. Agencies must also determine whether the current implementation of the memorandum's minimum practices adequately mitigates new and existing risks, or whether

³⁹ For example, the AI-specific exploits outlined in the MITRE ATLAS framework, see <https://atlas.mitre.org/> and NIST's taxonomy for adversarial machine learning, see <https://csrc.nist.gov/pubs/ai/100/2/e2023/final>.

updated risk response options are required.⁴⁰ At a minimum, human review is required at least on an annual basis and after significant modifications to the AI or to the conditions or context in which the AI is used, and the review must include renewed testing for performance of the AI in a real-world context.⁴¹ Reviews must also include oversight and consideration by an appropriate internal agency authority not directly involved in the system's development or operation.

- F. **Mitigate emerging risks to rights and safety.** Upon identifying new or significantly altered risks to rights or safety through ongoing monitoring, periodic review, or other mechanisms, agencies must take steps to mitigate those risks, including, as appropriate, through updating the AI to reduce its risks or implementing procedural or manual mitigations, such as more stringent human intervention requirements. As significant modifications make the existing implementation of the other minimum practices in this section less effective, such as by making training or documentation inaccurate, agencies must update or repeat those practices, as appropriate. Where the AI's risks to rights or safety exceed an acceptable level and where mitigation strategies do not sufficiently reduce risk, agencies must stop using the AI as soon as is practicable.⁴²
- G. **Ensure adequate human training and assessment.** Agencies must ensure there is sufficient training, assessment, and oversight for operators of the AI to interpret and act on the AI's output, combat any human-machine teaming issues (such as automation bias), and ensure the human-based components of the system effectively manage risks from the use of AI. Training should be conducted on a periodic basis, determined by the agency, and should be specific to the AI product or service being operated and how it is being used.
- H. **Provide additional human oversight, intervention, and accountability as part of decisions or actions that could result in a significant impact on rights or safety.** Agencies must assess their rights-impacting and safety-impacting uses of AI to identify any decisions or actions in which the AI is not permitted to act without additional human oversight, intervention, and accountability. When immediate human intervention is not practicable for such an action or decision, agencies must ensure that the AI functionality has an appropriate fail-safe that minimizes the risk of significant harm.⁴³

⁴⁰ In some cases, this may require a program evaluation, as defined under requirements of the Foundations for Evidence-Based Policymaking Act of 2018, Pub. L. No. 115-435, to determine the extent to which the AI is advancing the agency's mission and objectives.

⁴¹ For customer-facing services, agencies should consider customer feedback in their human review criteria.

⁴² Agencies are responsible for determining how to safely decommission AI that was already in use at the time of this memorandum's release, without significant disruptions to essential government functions.

⁴³ For example, an AI-enabled safety mechanism may require an immediate and automated action to prevent a harm from occurring. It would not be practicable in this case to require human intervention to approve the activation of the safety mechanism. However, agencies must still determine the appropriate oversight and accountability processes for such a use of AI.

- I. **Provide public notice and plain-language documentation.** Agencies must ensure, to the extent consistent with applicable law and governmentwide guidance, including concerning protection of privacy and of sensitive law enforcement, national security, and other protected information, that the AI's entry in the use case inventory provides 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. Where people interact with a service relying on the AI and are likely to be impacted by the AI, agencies must also provide reasonable and timely notice⁴⁴ about the use of the AI and a means to directly access any public documentation about it in the use case inventory. Where agencies' use cases are not included in their public inventories, they may still be required to report relevant information to OMB and must ensure adequate transparency in their use of AI, as appropriate and consistent with applicable law.
- v. **Additional Minimum Practices for Rights-Impacting AI.**

No later than December 1, 2024, agencies must follow the above minimum practices for AI that is *either* safety-impacting *or* rights-impacting. In addition, no later than December 1, 2024, agencies must also follow these minimum practices *before* initiating use of new or existing rights-impacting AI:

 - A. **Identify and assess AI's impact on equity and fairness, and mitigate algorithmic discrimination when it is present.** Agencies must:
 1. Identify and document in their AI impact assessment when using data that contains information about a class protected by Federal nondiscrimination laws (e.g., race, age, etc.). Given the risks arising when AI may correlate demographic information with other types of information, agencies should also assess and document whether the AI model could foreseeably use other attributes as proxies for a protected characteristic and whether such use would significantly influence model performance;
 2. Assess the AI in a real-world context to determine whether the AI model results in significant disparities in the model's performance (e.g., accuracy, precision, reliability in predicting outcomes) across demographic groups;
 3. Mitigate disparities that lead to, or perpetuate, unlawful discrimination or harmful bias, or that decrease equity as a result of the government's use of the AI; and
 4. Consistent with applicable law, cease use of the AI for agency decision-making if the agency is unable to adequately mitigate any associated risk of unlawful discrimination against protected classes. Agencies should maintain appropriate documentation to accompany this decision-making, and should disclose it publicly to the extent consistent with applicable law and governmentwide policy.

⁴⁴ Wherever feasible, agencies should provide notice to a user before the AI takes an action that significantly impacts them.

B. Consult and incorporate feedback from affected communities and the public.

Consistent with applicable law and governmentwide guidance, agencies must consult affected communities, including underserved communities, and they must solicit public feedback, where appropriate, in the design, development, and use of the AI and use such feedback to inform agency decision-making regarding the AI. The consultation and feedback process must include seeking input on the agency's approach to implementing the minimum risk management practices established in Section 5(c) of this memorandum, such as applicable opt-out procedures. Agencies should consider and manage the risks of public consultation in contexts like fraud prevention and law enforcement investigations, where consulting with the targeted individual is impractical but consulting with a representative group may be appropriate.⁴⁵

Agencies are strongly encouraged to solicit feedback on an ongoing basis from affected communities in particular as well as from the public broadly, especially after significant modifications to the AI or the conditions or context in which it is used.⁴⁶ In the course of assessing such feedback, if an agency determines that the use of AI in a given context would cause more harm than good, the agency should not use the AI.

To carry out such consultations and feedback processes, agencies must take appropriate steps to solicit input from the communities and individuals affected by the AI, which could include:⁴⁷

1. direct usability testing, such as observing users interacting with the system;
2. general solicitations of comments from the public, such as a request for information in the *Federal Register* or a "Tell Us About Your Experience" sheet with an open-ended space for responses;
3. post-transaction customer feedback collections;⁴⁸
4. public hearings or meetings, such as a listening session;
5. outreach to relevant Federal employee groups and Federal labor organizations, including on the appropriate fulfillment of collective bargaining obligations, where applicable; or
6. any other transparent process that seeks public input, comments, or feedback from the affected groups in a meaningful, equitable, accessible,

⁴⁵ For example, an agency using an AI tool to detect Federal benefits fraud is not required to consult with the target of their investigation. However, an agency should discern when it is appropriate to consult with civil society groups, academia, or other experts in the field to understand the technology's impact.

⁴⁶ The affected communities will vary depending on an agency's deployment context, but may include customers (for example, individuals, businesses, or organizations that interact with an agency) or Federal employee groups and employees' union representatives, when applicable.

⁴⁷ Agencies are encouraged to engage with OMB on whether they are required to submit information collection requests for OMB clearance under the Paperwork Reduction Act (44 U.S.C. § 3507) for the purposes of these consultations and feedback processes.

⁴⁸ Information on post-transaction customer feedback surveys can be found in OMB Circular A-11, Section 280 – Managing Customer Experience and Improving Service Delivery, <https://www.whitehouse.gov/wp-content/uploads/2018/06/s280.pdf>.

and effective manner.

No later than December 1, 2024 and on an ongoing basis *while* using new or existing covered rights-impacting AI, agencies must ensure these practices are followed for the AI:

- C. **Conduct ongoing monitoring and mitigation for AI-enabled discrimination.** As part of the ongoing monitoring requirement established in Section 5(c)(iv)(D), agencies must also monitor rights-impacting AI to specifically assess and mitigate AI-enabled discrimination against protected classes, including discrimination that might arise from unforeseen circumstances, changes to the system after deployment, or changes to the context of use or associated data. Where sufficient mitigation is not possible, agencies must safely discontinue use of the AI functionality.
- D. **Notify negatively affected individuals.** Consistent with applicable law and governmentwide guidance, agencies must notify individuals when use of the AI results in an adverse decision or action that specifically concerns them, such as the denial of benefits or deeming a transaction fraudulent.⁴⁹ Agencies should consider the timing of their notice and when it is appropriate to provide notice in multiple languages and through alternative formats and channels, depending on the context of the AI's use. The notice must also include a clear and accessible means of contacting the agency and, where applicable, provide information to the individual on their right to appeal. Agencies must also abide by any existing obligations to provide explanations for such decisions and actions.⁵⁰
- E. **Maintain human consideration and remedy processes.** Where practicable and consistent with applicable law and governmentwide guidance, agencies must provide timely human consideration and potential remedy, if appropriate, to the use of the AI via a fallback and escalation system in the event that an impacted individual would like to appeal or contest the AI's negative impacts on them. Agencies that already maintain an appeal or secondary human review process for adverse actions, or for agency officials' substantive or procedural errors, can leverage and expand such processes, as appropriate, or establish new processes to meet this requirement. These remedy processes should not place unnecessary burden on the impacted individual, and agencies should follow OMB guidance on

⁴⁹ In some instances, such as an active law enforcement investigation, providing immediate notice may be inappropriate or impractical, or disclosure may be more appropriate at a later stage (for example, prior to a defendant's trial).

⁵⁰ Explanations might include, for example, how and why the AI-driven decision or action was taken. This does not mean that agencies must provide a perfect breakdown of how a machine learning system came to a conclusion, as exact explanations of AI decisions may not be technically feasible. However, agencies should still characterize the general nature of such AI decisions through context such as the data that the decision relied upon, the design of the AI, and the broader decision-making context in which the system operates. Such explanations should be technologically valid, meaningful, useful, and as simply stated as possible, and higher-risk decisions should be accompanied by more comprehensive explanations.

calculating administrative burden.⁵¹ Whenever agencies are unable to provide an opportunity for an individual to appeal due to law, governmentwide guidance, or impracticability, they must create appropriate alternative mechanisms for human oversight of the AI.

- F. **Maintain options to opt-out for AI-enabled decisions.** Agencies must provide and maintain a mechanism for individuals to conveniently opt-out from the AI functionality in favor of a human alternative, where practicable and consistent with applicable law and governmentwide guidance. An opt-out mechanism must be prominent, readily available, and accessible, and it is especially critical where the affected people have a reasonable expectation of an alternative or where lack of an alternative would meaningfully limit availability of a service or create unwarranted harmful impacts. Agencies should also seek to ensure that the opt-out mechanism itself does not impose discriminatory burdens on access to a government service. Agencies are not required to provide the ability to opt-out if the AI functionality is solely used for the prevention, detection, and investigation of fraud⁵² or cybersecurity incidents, or the conduct of a criminal investigation. Pursuant to the authority for waivers defined in Section 5(c)(ii), CAIOs are additionally permitted to waive this opt-out requirement if they can demonstrate that a human alternative would result in a service that is less fair (e.g., produces a disparate impact on protected classes) or if an opt-out would impose undue hardship on the agency.

d. Managing Risks in Federal Procurement of Artificial Intelligence

This section provides agencies with recommendations for responsible procurement of AI, supplementing an agency's required risk management practices above for rights-impacting AI and safety-impacting AI. In addition to these recommendations and consistent with section 7224(d) of the Advancing American AI Act and Section 10.1(d)(ii) of Executive Order 14110, OMB will also develop an initial means to ensure that Federal contracts for the acquisition of an AI system or service align with the guidance in this memorandum.

- i. **Aligning with the Law.** Agencies should ensure that procured AI is consistent with the Constitution and complies with all other applicable laws, regulations, and policies, including those addressing privacy, confidentiality, intellectual property, cybersecurity, human and civil rights, and civil liberties.
- ii. **Transparency and Performance Improvement.** Agencies should take steps to ensure transparency and adequate performance for their procured AI, including by:
 - A. obtaining adequate documentation to assess the AI's capabilities, such as through the use of model, data, and system cards;

⁵¹ See OMB [M-22-10](#) and supporting document "[Strategies for Reducing Administrative Burden in Public Benefit and Service Programs](#)."

⁵² Some uses of AI in these categories, such as the use of biometrics for identity verification, may be subject to requirements in other guidance that would necessitate an option to opt-out, and this memorandum does not replace, supersede, otherwise interfere with any such requirements.

- B. obtaining adequate documentation of known limitations of the AI and any guidelines on how the system is intended to be used;
 - C. obtaining adequate information about the provenance of the data used to train, fine-tune, or operate the AI;
 - D. regularly evaluating claims made by Federal contractors concerning both the effectiveness of their AI offerings as well as the risk management measures put in place, including by testing the AI in the particular environment where the agency expects to deploy the capability;
 - E. considering contracting provisions that incentivize the continuous improvement of procured AI; and
 - F. requiring sufficient post-award monitoring of the AI, where appropriate in the context of the product or service acquired.
- iii. **Promoting Competition in Procurement of AI.** Agencies should take appropriate steps to ensure that Federal AI procurement practices promote opportunities for competition among contractors and do not improperly entrench incumbents. Such steps may include promoting interoperability so that, for example, procured AI works across multiple cloud environments, and ensuring that vendors do not inappropriately favor their own products at the expense of competitors' offerings.
- iv. **Maximizing the Value of Data for AI.** In contracts for AI products and services, agencies should treat relevant data, as well as improvements to that data—such as cleaning and labeling—as a critical asset for their AI maturity. Agencies should take steps to ensure that their contracts retain for the Government sufficient rights to data and any improvements to that data so as to avoid vendor lock-in and facilitate the Government's continued design, development, testing, and operation of AI. Additionally, agencies should consider contracting provisions that protect Federal information used by vendors in the development and operation of AI products and services for the Federal Government, so that such data is protected from unauthorized disclosure and use and cannot be subsequently used to train or improve the functionality of the vendor's commercial offerings without express permission from the agency.
- v. **Overfitting to Known Test Data.** When testing AI using data that its developer may have access to—including test data that the agency has itself shared or released—agencies should ensure, as appropriate, that their AI developers or vendors are not directly relying on the test data to train their AI systems.⁵³
- vi. **Responsible Procurement of AI for Biometric Identification.** When procuring systems that use AI to identify individuals using biometric identifiers—e.g., faces, irises, fingerprints, or gait—agencies are encouraged to:
 - A. Assess and address the risks that the data used to train or operate the AI may not be lawfully collected or used, or else may not be sufficiently accurate to support reliable biometric identification. This includes the risks that the biometric information was collected without appropriate consent, was originally collected

⁵³ For instance, using validation data to train a model could lead the model to learn spurious correlations that make the model appear accurate in tests but harm the real-world performance of the AI system.

for another purpose, embeds unwanted bias, or was collected without validation of the included identities; and

- B. Request supporting documentation or test results to validate the accuracy, reliability, and validity of the AI's ability to match identities.

vii. **Responsibly Procuring Generative AI.** Agencies are encouraged to include risk management requirements in contracts for generative AI, and particularly for dual-use foundation models, including:

- A. requiring adequate testing and safeguards,
- B. requiring results of internal or external testing and evaluation, to include AI red-teaming against risks from generative AI, such as discriminatory, misleading, inflammatory, unsafe, or deceptive outputs;
- C. requiring that generative AI models have capabilities, as appropriate and technologically feasible, to reliably label or establish provenance for their content as generated or modified by AI; and
- D. incorporating relevant NIST standards, defined pursuant to Sections 4.1(a) and 10.1(d) of Executive Order 14110, as appropriate.

viii. **Assessing for Environmental Efficiency and Sustainability.** When procuring computationally intensive AI services, for example those that rely on dual-use foundation models, agencies should consider the environmental impact of those services, including whether the vendor has implemented methods to improve the efficiency and sustainability of such AI. This should include considering the carbon emissions and resource consumption from supporting data centers.

6. DEFINITIONS

The below definitions apply for the purposes of this memorandum.

Accessibility: The term “accessibility” has the meaning provided in Section 2(e) of Executive Order 14035.

Agency: The term “agency” has the meaning provided in 44 U.S.C. § 3502(1).

Algorithmic Discrimination: The term “algorithmic discrimination” has the meaning provided in Section 10(f) of Executive Order 14091 of February 16, 2023.

Artificial Intelligence (AI): The term “artificial intelligence” has the meaning provided in Section 238(g) of the John S. McCain National Defense Authorization Act for Fiscal Year 2019,⁵⁴ which states that “the term ‘artificial intelligence’ includes the following”:

1. Any artificial system that performs tasks under varying and unpredictable circumstances without significant human oversight, or that can learn from experience and improve performance when exposed to data sets.

⁵⁴ Pub. L. No. 115-232, § 238(g), <https://www.govinfo.gov/content/pkg/PLAW-115publ232/pdf/PLAW-115publ232.pdf>.

2. An artificial system developed in computer software, physical hardware, or other context that solves tasks requiring human-like perception, cognition, planning, learning, communication, or physical action.
3. An artificial system designed to think or act like a human, including cognitive architectures and neural networks.
4. A set of techniques, including machine learning, that is designed to approximate a cognitive task.
5. An artificial system designed to act rationally, including an intelligent software agent or embodied robot that achieves goals using perception, planning, reasoning, learning, communicating, decision making, and acting.

For the purposes of this memorandum, the following technical context should guide interpretation of the definition above:

1. This definition of AI encompasses, but is not limited to, the AI technical subfields of machine learning (including deep learning as well as supervised, unsupervised, and semi-supervised approaches), reinforcement learning, transfer learning, and generative AI.
2. This definition of AI does not include robotic process automation or other systems whose behavior is defined only by human-defined rules or that learn solely by repeating an observed practice exactly as it was conducted.
3. For this definition, no system should be considered too simple to qualify as covered AI due to a lack of technical complexity (e.g., the smaller number of parameters in a model, the type of model, or the amount of data used for training purposes).
4. This definition includes systems that are fully autonomous, partially autonomous, and not autonomous, and it includes systems that operate both with and without human oversight.

AI and AI-Enabling Roles: The term “AI and AI-enabling roles” refers to individuals with positions and major duties whose contributions are important for successful and responsible AI outcomes. AI and AI-Enabling Roles include both technical and non-technical roles, such as data scientists, software engineers, data engineers, data governance specialists, statisticians, machine learning engineers, applied scientists, designers, economists, operations researchers, product managers, policy analysts, program managers, behavioral and social scientists, customer experience strategists, human resource specialists, contracting officials, managers, and attorneys.

AI Maturity: The term “AI maturity” refers to a Federal Government organization’s capacity to successfully and responsibly adopt AI into their operations and decision-making across the organization, manage its risks, and comply with relevant Federal law, regulation, and policy on AI.

AI Model: The term “AI model” has the meaning provided in Section 3(c) of Executive Order 14110.

AI Red-Teaming: The term “AI red-teaming” has the meaning provided for “AI red-teaming” in Section 3(d) of Executive Order 14110.

Applied Research: The term “applied research” refers to original investigation undertaken in order to acquire new knowledge to determine the means by which a specific practical aim or objective may be met.

Automation Bias: The term “automation bias” refers to the propensity for humans to inordinately favor suggestions from automated decision-making systems and to ignore or fail to seek out contradictory information made without automation.

Basic Research: The term “basic research” refers to experimental or theoretical work undertaken primarily to acquire new knowledge of the underlying foundations of phenomena and observable facts without a specific application towards processes or products in mind.

CFO Act Agency: The term “CFO Act Agency” refers to the agencies identified in 31 U.S.C. § 901(b).

Custom-Developed Code: The term “custom-developed code” has the meaning provided in Appendix A of OMB Memorandum M-16-21.

Customer Experience: The term “customer experience” has the meaning established in Section 3(b) of Executive Order 14058.⁵⁵

Data Asset: The term “data asset” has the meaning provided in 44 U.S.C § 3502.

Dual-Use Foundation Model: The term “dual-use foundation model” has the meaning provided in Section 3(k) of Executive Order 14110.

Equity: The term “equity” has the meaning provided in Section 10(a) of Executive Order 14091.⁵⁶

Federal Information: The term “Federal information” has the meaning provided in OMB Circular A-130.

Generative AI: The term “generative AI” has the meaning provided in Section 3(p) of Executive Order 14110.

Intelligence Community: The term “intelligence community” has the meaning provided in 50 U.S.C. § 3003.

Model Weight: The term “model weight” has the meaning provided in Section 3(u) of Executive Order 14110.

⁵⁵ Executive Order 14058, *Transforming Federal Customer Experience and Service Delivery To Rebuild Trust in Government*, <https://www.federalregister.gov/documents/2021/12/16/2021-27380/transforming-federal-customer-experience-and-service-delivery-to-rebuild-trust-in-government>.

⁵⁶ Executive Order 14091, *Further Advancing Racial Equity and Support for Underserved Communities Through the Federal Government*, <https://www.govinfo.gov/content/pkg/FR-2023-02-22/pdf/2023-03779.pdf>.

National Security System: The term “National Security System” has the meaning provided in 44 U.S.C. § 3552(b)(6).

Open Government Data Asset: The term “open government data asset” has the meaning provided in 44 U.S.C § 3502.

Open Source Software: The term “open source software” has the meaning provided in Appendix A of OMB Memorandum M-16-21.

Rights-Impacting AI:⁵⁷ The term “rights-impacting AI” refers to AI whose output serves as a principal basis for a decision or action concerning a specific individual or entity that has a legal, material, binding, or similarly significant effect on that individual’s or entity’s:

1. Civil rights, civil liberties, or privacy, including but not limited to freedom of speech, voting, human autonomy, and protections from discrimination, excessive punishment, and unlawful surveillance;
2. Equal opportunities, including equitable access to education, housing, insurance, credit, employment, and other programs where civil rights and equal opportunity protections apply; or
3. Access to or the ability to apply for critical government resources or services, including healthcare, financial services, public housing, social services, transportation, and essential goods and services.

Risks from the Use of AI: The term “risks from the use of AI” refers to risks related to efficacy, safety, equity, fairness, transparency, accountability, appropriateness, or lawfulness of a decision or action resulting from the use of AI to inform, influence, decide, or execute that decision or action. This includes such risks regardless of whether:

1. the AI merely informs the decision or action, partially automates it, or fully automates it;
2. there is or is not human oversight for the decision or action;
3. it is or is not easily apparent that a decision or action took place, such as when an AI application performs a background task or silently declines to take an action; or
4. the humans involved in making the decision or action or that are affected by it are or are not aware of how or to what extent the AI influenced or automated the decision or action.

While the particular forms of these risks continue to evolve, at least the following factors can create, contribute to, or exacerbate these risks:

1. AI outputs that are inaccurate or misleading;
2. AI outputs that are unreliable, ineffective, or not robust;
3. AI outputs that are discriminatory or have a discriminatory effect;
4. AI outputs that contribute to actions or decisions resulting in harmful or unsafe outcomes, including AI outputs that lower the barrier for people to take intentional and harmful actions;
5. AI being used for tasks to which it is poorly suited or being inappropriately repurposed in a context for which it was not intended;
6. AI being used in a context in which affected people have a reasonable expectation that a human is or should be primarily responsible for a decision or action; and

⁵⁷ Appendix I(2) of this memorandum lists AI applications that are presumed to be rights-impacting.

7. the adversarial evasion or manipulation of AI, such as an entity purposefully inducing AI to misclassify an input.

This definition applies to risks specifically arising from using AI and that affect the outcomes of decisions or actions. It does not include all risks associated with AI, such as risks related to the privacy, security, and confidentiality of the data used to train AI or used as inputs to AI models.

Safety-Impacting AI:⁵⁸ The term “safety-impacting AI” refers to AI whose output produces an action or serves as a principal basis for a decision that has the potential to significantly impact the safety of:

1. Human life or well-being, including loss of life, serious injury, bodily harm, biological or chemical harms, occupational hazards, harassment or abuse, or mental health, including both individual and community aspects of these harms;
2. Climate or environment, including irreversible or significant environmental damage;
3. Critical infrastructure, including the critical infrastructure sectors defined in Presidential Policy Directive 21⁵⁹ or any successor directive and the infrastructure for voting and protecting the integrity of elections; or,
4. Strategic assets or resources, including high-value property and information marked as sensitive or classified by the Federal Government.

Significant Modification: The term “significant modification” refers to an update to an AI application or to the conditions or context in which it is used that meaningfully alters the AI’s impact on rights or safety, such as through changing its functionality, underlying structure, or performance such that prior evaluations, training, or documentation become misleading to users, overseers, or individuals affected by the system. This includes significantly changing the context, scope, or intended purpose in which the AI is used.

Underserved Communities: The term “underserved communities” has the meaning provided in Section 10(b) of Executive Order 14091.

⁵⁸ Appendix I(1) of this memorandum lists AI applications that are presumed to be safety-impacting.

⁵⁹ Presidential Policy Directive 21 (PPD-21), *Critical Infrastructure Security and Resilience*, or successor directive, <https://obamawhitehouse.archives.gov/the-press-office/2013/02/12/presidential-policy-directive-critical-infrastructure-security-and-resil>.

Appendix I: Purposes for Which AI is Presumed to be Safety-Impacting and Rights-Impacting

OMB has determined that the categories in this appendix in general meet the definition of safety-impacting AI or rights-impacting AI and are automatically *presumed* to be safety-impacting or rights-impacting. The following lists only identify a subset of uses of AI that impact rights and safety, and they do not represent an exhaustive list. Additionally, the presumption that a particular use of AI in the following lists will impact rights or safety can be waived by an agency's CAIO with adequate justification, pursuant to the processes outlined in Section 5.

1. Purposes That Are Presumed to Be Safety-Impacting. A use of AI is presumed to be safety-impacting if it is used or expected to be used, in real-world conditions, to control or significantly influence the outcomes of any of the following agency activities or decisions:

- a. Controlling the safety-critical functions within dams, emergency services, electrical grids, the generation or movement of energy, fire safety systems, food safety mechanisms, traffic control systems and other systems controlling physical transit, water and wastewater systems, or nuclear reactors, materials, and waste;
- b. Maintaining the integrity of elections and voting infrastructure;
- c. Controlling the physical movements of robots or robotic appendages within a workplace, school, housing, transportation, medical, or law enforcement setting;
- d. Applying kinetic force; delivering biological or chemical agents; or delivering potentially damaging electromagnetic impulses;
- e. Autonomously or semi-autonomously moving vehicles, whether on land, underground, at sea, in the air, or in space;
- f. Controlling the transport, safety, design, or development of hazardous chemicals or biological agents;
- g. Controlling industrial emissions and environmental impacts;
- h. Transporting or managing of industrial waste or other controlled pollutants;
- i. Designing, constructing, or testing of industrial equipment, systems, or structures that, if they failed, would pose a significant risk to safety;
- j. Carrying out the medically relevant functions of medical devices; providing medical diagnoses; determining medical treatments; providing medical or insurance health-risk assessments; providing drug-addiction risk assessments or determining access to medication; conducting risk assessments for suicide or other violence; detecting or preventing mental-health issues; flagging patients for interventions; allocating care in the context of public insurance; or controlling health-insurance costs and underwriting;
- k. Detecting the presence of dangerous weapons or a violent act;
- l. Choosing to summon first responders to an emergency;
- m. Controlling access to or security of government facilities; or
- n. Determining or carrying out enforcement actions pursuant to sanctions, trade restrictions, or other controls on exports, investments, or shipping.

2. Purposes That Are Presumed to Be Rights-Impacting. A use of AI is presumed to be rights-impacting if it is used or expected to be used, in real-world conditions, to control or significantly influence the outcomes of any of the following agency activities or decisions:

- a. Blocking, removing, hiding, or limiting the reach of protected speech;
- b. In law enforcement contexts, producing risk assessments about individuals; predicting criminal recidivism; predicting criminal offenders; identifying criminal suspects or predicting perpetrators' identities; predicting victims of crime; forecasting crime; detecting gunshots; tracking personal vehicles over time in public spaces, including license plate readers; conducting biometric identification (e.g., iris, facial, fingerprint, or gait matching); sketching faces; reconstructing faces based on genetic information; monitoring social media; monitoring prisons; forensically analyzing criminal evidence; conducting forensic genetics; conducting cyber intrusions in the course of an investigation; conducting physical location-monitoring or tracking of individuals; or making determinations related to sentencing, parole, supervised release, probation, bail, pretrial release, or pretrial detention;
- c. Deciding or providing risk assessments related to immigration, asylum, or detention status; providing immigration-related risk assessments about individuals who intend to travel to, or have already entered, the U.S. or its territories; determining individuals' border access or access to Federal immigration related services through biometrics or through monitoring social media and other online activity; monitoring individuals' physical location for immigration and detention-related purposes; or forecasting the migration activity of individuals;
- d. Conducting biometric identification for one-to-many identification in publicly accessible spaces;
- e. Detecting or measuring emotions, thought, impairment, or deception in humans;
- f. Replicating a person's likeness or voice without express consent;
- g. In education contexts, detecting student cheating or plagiarism; influencing admissions processes; monitoring students online or in virtual-reality; projecting student progress or outcomes; recommending disciplinary interventions; determining access to educational resources or programs; determining eligibility for student aid or Federal education; or facilitating surveillance (whether online or in-person);
- h. Screening tenants; monitoring tenants in the context of public housing; providing valuations for homes; underwriting mortgages; or determining access to or terms of home insurance;
- i. Determining the terms or conditions of employment, including pre-employment screening, reasonable accommodation, pay or promotion, performance management, hiring or termination, or recommending disciplinary action; performing time-on-task tracking; or conducting workplace surveillance or automated personnel management;
- j. Carrying out the medically relevant functions of medical devices; providing medical diagnoses; determining medical treatments; providing medical or insurance health-risk assessments; providing drug-addiction risk assessments or determining access to medication; conducting risk assessments for suicide or other violence; detecting or preventing mental-health issues; flagging patients for interventions; allocating care in the context of public insurance; or controlling health-insurance costs and underwriting;

- k. Allocating loans; determining financial-system access; credit scoring; determining who is subject to a financial audit; making insurance determinations and risk assessments; determining interest rates; or determining financial penalties (e.g., garnishing wages or withholding tax returns);
- l. Making decisions regarding access to, eligibility for, or revocation of critical government resources or services; allowing or denying access—through biometrics or other means (e.g., signature matching)—to IT systems for accessing services for benefits; detecting fraudulent use or attempted use of government services; assigning penalties in the context of government benefits;
- m. Translating between languages for the purpose of official communication to an individual where the responses are legally binding; providing live language interpretation or translation, without a competent interpreter or translator present, for an interaction that directly informs an agency decision or action; or
- n. Providing recommendations, decisions, or risk assessments about adoption matching, child protective actions, recommending child custody, whether a parent or guardian is suitable to gain or retain custody of a child, or protective actions for senior citizens or disabled persons.

Appendix II: Consolidated Table of Actions

Responsible Entity	Action	Section	Deadline
Each Agency	Designate an agency Chief AI Officer and notify OMB	3(a)(i)	60 days
Each CFO Act Agency	Convene agency AI Governance Board	3(a)(ii)	60 days
Each Agency	Submit to OMB and release publicly an agency plan to achieve consistency with this memorandum or a written determination that the agency does not use and does not anticipate using covered AI	3(a)(iii)	180 days and every two years thereafter until 2036
Each CFO Act Agency	Develop and release publicly an agency strategy for removing barriers to the use of AI and advancing agency AI maturity	4(a)(i)	365 days
Each Agency**	Publicly release an expanded AI use case inventory and report metrics on use cases not included in public inventories	3(a)(iv), 3(a)(v)	Annually
Each Agency*	Share and release AI code, models, and data assets, as appropriate	4(d)	Ongoing
Each Agency*	Stop using any safety-impacting or rights-impacting AI that is not in compliance with Section 5(c) and has not received an extension or waiver	5(a)(i)	December 1, 2024 (with extensions possible)
Each Agency*	Certify the ongoing validity of the waivers and determinations granted under Section 5(c) and 5(b) and publicly release a summary detailing each and its justification	5(a)(ii)	December 1, 2024 and annually thereafter
Each Agency*	Conduct periodic risk reviews of any safety-impacting and rights-impacting AI in use	5(c)(iv)(D)	At least annually and after significant modifications
Each Agency*	Report to OMB any determinations made under Section 5(b) or waivers granted under Section 5(c)	5(b); 5(c)(iii)	Ongoing, within 30 days of granting waiver

* Excluding elements of the Intelligence Community.

** Excluding elements of the Intelligence Community. The Department of Defense is exempt from the requirement to inventory individual use cases.