



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

THE DIRECTOR

PROPOSED 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 Artificial Intelligence in Government Act of 2020,¹ the Advancing American AI Act,² and President Biden’s Executive Order of October 30, 2023 (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, particularly those affecting the safety and rights of the public.

As set forth in the accompanying Federal Register notice, the Office of Management and Budget is requesting public comment on this proposed memorandum.

1. OVERVIEW

While AI is improving operations and efficiency 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 President Biden’s October 30, 2023 Executive Order (the “AI Executive Order”), 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. Because AI is deeply interconnected with other technical and policy areas including data, information

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>.

technology (IT), security, privacy, civil rights and civil liberties, customer experience, and workforce management, CAIOs must work in close coordination with existing responsible officials and organizations within their agencies.

Advancing Responsible AI Innovation. When implemented responsibly, AI can improve operations across the Federal Government. Agencies must increase their capacity to successfully and responsibly adopt AI, including generative AI, into their operations. To that end, this memorandum requires each agency identified in the Chief Financial Officer (CFO) Act³ to develop an enterprise AI strategy. 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 creates new requirements focused specifically on the 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.⁴ To address these risks, this memorandum requires agencies to follow minimum practices when using rights-impacting and safety-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.

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 scoped to address risks specifically arising from the use of AI, as well as governance and innovation issues that are directly tied to agencies' use of AI. This memorandum does not address issues that are present regardless of the use of AI, for instance with respect to Federal information and 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, 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.

³ 31 U.S.C. § 901(b).

⁴ A full definition for "risks from the use of AI" is provided in Section 6.

a. Covered Agencies. Except as specifically noted, this memorandum applies to all agencies defined in 44 U.S.C. § 3502(1).⁵ As noted in the relevant sections, some requirements in this memorandum only apply to 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. The principles of this memorandum do not, by contrast, govern agencies' regulatory actions designed to prescribe law or policy regarding non-agency uses of AI.

The requirements 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 used as a component of a national security system.⁶

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 that risks from the agency's use of AI are adequately managed. The head of each covered agency must also consider the necessary financial, human, information, and infrastructural resources to carry out these responsibilities effectively, including providing or requesting resources via the budget process 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 the AI Executive Order. CAIOs bear primary responsibility on

⁵ The term "agency" 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. As a result, agencies defined in 44 U.S.C. § 3502(5) (independent regulatory agencies) that were not covered by Executive Order 13960 of December 8, 2020 *are* covered by this memorandum.

⁶ 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 the AI Executive Order requires the development of a National Security Memorandum to govern the use of AI as a component of a National Security System, and agencies 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*.

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, and they must update OMB within 30 days when the designated individual changes.
- 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 coordinate and govern AI issues, consistent with Section 10.1(b) of the AI Executive Order and the detailed guidance in Section 3(c) of this memorandum.
- iii. **Compliance Plans.** Consistent with Section 104(c)-(d) of the AI in Government Act, 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.⁷ OMB will provide full templates for these compliance plans.
- iv. **AI Use Case Inventories.** Pursuant to Section 7225 of the Advancing American AI Act, and subject to the exclusions in that Act and Section 10.1(e) of the AI Executive Order, each agency (except for the Department of Defense and the Intelligence Community) must annually submit an inventory of its AI use cases to OMB and subsequently post a public version on the agency's website.⁸ OMB will issue detailed instructions for the inventory 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 and report additional detail on how they are using safety-impacting and rights-impacting AI, the risks—including risks to equity—that such use poses, how they are managing those risks, and any related extensions and waivers granted under

⁷ 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 the guidance in this memorandum.

⁸ Agencies must only publicly report use cases to the extent practicable and consistent with applicable law and governmentwide guidance, including those concerning the protection of privacy and of sensitive law enforcement, national security, and other protected information.

Section 5 of this memorandum.

- v. **Reporting on AI Use Cases Not Subject to Inventory.** Some AI use cases are exempt from the Advancing American AI Act’s inventory requirement. Of those use cases, those within the Department of Defense are otherwise within the scope of this memorandum unless they concern AI used as a component of a national security system. The Department of Defense must annually provide OMB with information on those in-scope AI use cases, including aggregate metrics about those in-scope AI uses cases, the number of such cases that impact rights and safety, their compliance with the practices of Section 5(c) of this memorandum, and any waivers granted under Section 5 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 the AI Executive Order, this memorandum defines agency 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, the CAIO’s primary role must be coordination, innovation, and risk management for their agency’s use of AI. Agencies may choose to designate an existing official, such as a Chief Technology Officer, Chief Data Officer, or similar official with relevant or complementary authorities and responsibilities, provided they have significant expertise in AI and meet the other requirements in this section.
- ii. **Responsibilities.** The AI Executive Order 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 13960⁹ and Section 4(b) of Executive Order 14091.¹⁰ 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. maintaining awareness of agency AI activities, including through creating and maintaining the annual AI use case inventory;

⁹ 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>.

¹⁰ 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>.

- C. 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;
- D. advising the agency CFO and Chief Human Capital Officer (CHCO) on the resourcing requirements and workforce skillsets necessary for applying AI to the agency's mission and adequately managing its risks;
- E. supporting agency involvement with appropriate interagency coordination bodies related to their agency's AI activities, including representing the agency to the council described in Section 10.1(a) of the AI Executive Order;
- F. 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;¹¹

Promoting AI Innovation

- G. working with their agency to identify and prioritize appropriate uses of AI that will improve their agency's mission and advance equity;
- H. identifying and removing barriers to the responsible use of AI in the agency, including through the advancement of AI-enabling enterprise infrastructure, workforce development measures, policy, and other resources for AI innovation;
- I. 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

- J. 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;
- K. working with relevant senior agency officials to establish or update processes to measure, monitor, and evaluate the ongoing performance of AI applications and whether they are achieving their intended objectives;
- L. 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;
- M. conducting risk assessments, as necessary, of agency AI applications to ensure compliance with this memorandum;
- N. overseeing development of agency-specific lists, as necessary, of purposes for which AI is presumed to be safety-impacting or rights-impacting;¹²
- O. waiving individual applications of AI from elements of Section 5 of this memorandum through the processes detailed in that section; and

¹¹ 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>.

¹² See Section 5(b) of this memorandum for the OMB-defined lists to which agency-specific lists would add. Any agency-specific lists will be governed by the same processes defined in Section 5(b) for the OMB-defined lists.

- P. in partnership with relevant agency officials (e.g., authorizing, procurement, legal, human capital, and oversight officials), ensuring 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 the AI Executive Order, 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,¹³ including in the development of the agency risk management strategy.¹⁴ 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.

CFO Act agencies are required specifically to establish AI Governance Boards to convene relevant senior officials no less than quarterly 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 bodies¹⁵ to fulfill this requirement as long as they currently satisfy or are made to satisfy both of the following:

¹³ 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.

¹⁴ 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.

¹⁵ An example of a qualifying body includes agency Data Governance Bodies, established by OMB Memorandum M-19-23, *Phase 1 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>.

- 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 responsibilities should not be assigned to other officials. Working through this Board, CAIOs will support their respective Deputy Secretaries in coordinating AI activities across the agency and implementing relevant sections of the AI Executive Order.
- 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, human capital, procurement, budget, agency management, customer experience, performance evaluation, statistics, risk management, equity, privacy, civil rights and civil liberties, 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.

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 of AI and when it is not suited for a given task. To achieve this, agencies should build internal enterprise capacity to support responsible AI innovation and take actions to improve their procurement of AI.

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 advances in AI maturity, including:

- i. the agency's current and planned top use cases of AI¹⁶;
- ii. a current assessment of the agency's AI maturity and the agency's AI maturity goals based on the method established under Section 10.1(c) of the AI Executive Order;
- iii. the agency's plans to effectively govern its use of AI, including through its Chief AI Officer, AI Governance Boards, and improvements to their 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 building sufficient enterprise capacity to manage risks from the use of AI;
- vi. a current assessment of the agency's AI workforce capacity and projected AI 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; and

¹⁶ 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 the Intelligence Community, as defined in 5 U.S.C. § 3003(4).

- vii. specific, prioritized areas and planning for future AI investment.

b. Removing Barriers to the Responsible Use of Artificial Intelligence

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 flexibility and do not face hindrances that divert limited resources and expertise away from the AI innovation and risk management. Agencies should take steps to remove such barriers, 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 monitoring capabilities necessary to rapidly develop, test, and maintain AI applications.
- ii. **Data.** Agencies should develop adequate infrastructure and capacity to sufficiently curate agency datasets for use in training, testing, and operating AI. This includes an agency's capacity to maximize appropriate access to internal data and share such data within the agency. Agencies should also explore the utility of public access datasets and encourage their use, where appropriate and consistent with the data practices outlined in this memorandum, to help develop, test, and maintain AI applications. These activities should be supported by resources to enable sound data governance and management practices, particularly as it relates to data curation, labeling, and stewardship.
- iii. **Cybersecurity.** Agencies should update, as necessary, cybersecurity authorization processes 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 the AI Executive Order, agency authorizing officials should also prioritize generative AI and other critical emerging technologies in Authorizations to Operate and any other applicable release or oversight processes.
- iv. **Workforce.** Consistent with Sections 5.1 and 10.2 of the AI Executive Order, agencies should take full advantage of available special hiring and retention authorities to fill gaps in AI talent, encouraging applications from individuals with diverse perspectives and experiences, and ensure the use of recruitment best practices for AI positions, such as descriptive job titles and skills-based assessments. When identifying and filling workforce needs for AI, agencies should include both technical roles, such as data scientists and engineers, and non-technical roles, such as designers, behavioral scientists, contracting officials, managers, and attorneys, whose contribution and competence with AI are important for successful and responsible AI outcomes. Agencies should provide resources and training to develop such AI talent internally and should also 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.

- v. **Generative AI.** In addition to heeding the guidance provided in Section 10.1(f) of the AI Executive Order, agencies should assess potential beneficial use cases 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.

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 that are not elements of the Intelligence Community are required to implement minimum practices, detailed below, to manage risks from rights-impacting and safety-impacting AI.¹⁷

a. Actions

- i. **Implementation of Risk Management Practices and Termination of Non-Compliant AI.** By August 1, 2024, agencies must implement the minimum practices in Section 5(c) of this memorandum for safety-impacting or rights-impacting AI, or else stop using any AI that is not compliant with the minimum practices, consistent with the details and caveats in that section.
- ii. **Recommendation on AI Documentation.** Within 180 days of the issuance of this memorandum, the council described in Section 10.1(a) of the AI Executive Order will provide the Director of OMB with a list of recommended documentation that should be required from a selected vendor in the fulfillment of a Federal AI contract. As part of their recommendation, the council must consider the minimum risk management practices in Section 5(c) and the associated materials that may be required of vendors to demonstrate that they have completed such tasks.

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

All AI within the scope of this section 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 appropriate deadline. Agencies must review each use of AI that they are developing or using to determine whether it matches the definition of safety-impacting or rights-impacting.

The categories in this subsection only identify a subset of specific purposes for which AI is automatically *presumed* to be safety-impacting or rights-impacting, and they do not represent an exhaustive list of purposes for which AI is safety-impacting or rights-impacting. Agencies are

¹⁷ Although elements of the Intelligence Community are not required to implement these practices, they are encouraged to do so.

also encouraged to define specific purposes that, within their agency, are presumed to be safety-impacting or rights-impacting and so must follow the practices in Section 5(c). Agencies are required to report any such agency-specific lists to OMB on an annual basis.

Where an agency currently uses or plans to use AI for a purpose described below, the CAIO, in coordination with other relevant officials as specified by the agency, may make a determination (or reverse a prior determination) that the AI application or component¹⁸ 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 reverse this determination only with a documented context-specific and system-specific risk assessment. Any such determination or reversal must be reported to OMB within 30 days.

- i. **Purposes That Are Presumed to Be Safety-Impacting.** Unless the CAIO determines otherwise, covered AI within the scope of this memorandum is presumed to be safety-impacting and must follow the minimum practices for safety-impacting AI if it is used to control or meaningfully influence the outcomes of the following activities:
 - A. The functioning of dams, emergency services, electrical grids or the generation or movement of energy, fire safety systems, food safety mechanisms, integrity of elections and voting infrastructure, traffic control systems and other systems controlling physical transit, water and wastewater systems, and nuclear reactors, materials, and waste;
 - B. Physical movements, including in human-robot teaming, such as the movements of a robotic appendage or body, within a workplace, school, housing, transportation, medical, or law enforcement setting;
 - C. The application of kinetic force, delivery of biological or chemical agents, or delivery of potentially damaging electromagnetic impulses;
 - D. The movements of vehicles, whether on land, underground, at sea, in the air, or in space;
 - E. The transport, safety, design, or development of hazardous chemicals or biological entities or pathways;
 - F. Industrial emissions and environmental impact control processes;
 - G. The transportation or management of industrial waste or other controlled pollutants;
 - H. The design, construction, or testing of industrial equipment, systems, or structures that, if they failed, would pose a meaningful risk to safety;
 - I. Responses to insider threats;
 - J. Access to or security of government facilities; or
 - K. Enforcement actions pursuant to sanctions, trade restrictions, or other controls on exports, investments, or shipping.

¹⁸ CAIOs may also make these determinations across groups of closely related AI applications or components, provided that: (1) those systems have undergone a risk assessment that adequately considers the risks from each individual system; and (2) the systems are substantially identical in their risk profiles.

- ii. **Purposes That Are Presumed to Be Rights-Impacting.** Unless the CAIO determines otherwise, covered AI is presumed to be rights-impacting (and potentially also safety-impacting) and agencies must follow the minimum practices for rights-impacting AI and safety-impacting AI if it is used to control or meaningfully influence the outcomes of any of the following activities or decisions:
- A. Decisions to block, remove, hide, or limit the reach of protected speech;
 - B. Law enforcement or surveillance-related risk assessments about individuals, criminal recidivism prediction, offender prediction, predicting perpetrators' identities, victim prediction, crime forecasting, license plate readers, iris matching, facial matching, facial sketching, genetic facial reconstruction, social media monitoring, prison monitoring, forensic analysis, forensic genetics, the conduct of cyber intrusions, physical location-monitoring devices, or decisions related to sentencing, parole, supervised release, probation, bail, pretrial release, or pretrial detention;
 - C. Deciding immigration, asylum, or detention status; providing risk assessments about individuals who intend to travel to, or have already entered, the U.S. or its territories; determining border access or access to Federal immigration related services through biometrics (e.g., facial matching) or other means (e.g., monitoring of social media or protected online speech); translating official communication to an individual in an immigration, asylum, detention, or border context; or immigration, asylum, or detention-related physical location-monitoring devices.
 - D. Detecting or measuring emotions, thought, or deception in humans;
 - E. In education, 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 facilitating surveillance (whether online or in-person);
 - F. Tenant screening or controls, home valuation, mortgage underwriting, or determining access to or terms of home insurance;
 - G. Determining the terms and conditions of employment, including pre-employment screening, pay or promotion, performance management, hiring or termination, time-on-task tracking, virtual or augmented reality workplace training programs, or electronic workplace surveillance and management systems;
 - H. Decisions regarding medical devices, medical diagnostic tools, clinical diagnosis and determination of treatment, medical or insurance health-risk assessments, drug-addiction risk assessments and associated access systems, suicide or other violence risk assessment, mental-health status detection or prevention, systems that flag patients for interventions, public insurance care-allocation systems, or health-insurance cost and underwriting processes;
 - I. Loan-allocation processes, financial-system access determinations, credit scoring, determining who is subject to a financial audit, insurance processes including risk

assessments, interest rate determinations, or financial systems that apply penalties (e.g., that can garnish wages or withhold tax returns);

- J. Decisions regarding access to, eligibility for, or revocation of government benefits or services; allowing or denying access—through biometrics or other means (e.g., signature matching)—to IT systems for accessing services for benefits; detecting fraud; assigning penalties in the context of government benefits; or
- K. Recommendations or decisions about child welfare, child custody, or whether a parent or guardian is suitable to gain or retain custody of a child.

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 practices in this section to safety-impacting and rights-impacting AI by August 1, 2024, or else stop using the AI until it becomes compliant. Prior to August 1, 2024, agency CAIOs should work with their agencies' relevant officials to bring potentially non-compliant AI into conformity, which may include voluntary requests to third-party vendors to 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. An agency may also request an extension or grant a waiver to this requirement through its CAIO using the processes detailed below.

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 such as a TechStat process,¹⁹ or on request as determined by OMB.

The practices in this section represent a minimum baseline for managing risk from the use of AI. Agencies must identify additional context-specific risks that are associated with their determined use cases 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, potential harms to worker wellbeing, access to critical resources and services, and effects on market competition. To fill 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,²⁰ the Blueprint for an AI Bill of Rights,²¹ applicable international standards,²² and the workforce principles established pursuant to Section 6 of the AI Executive

¹⁹ *Policies & Initiatives: TechStat*, U.S. Chief Information Officers Council, <https://www.cio.gov/handbook/policies-initiatives/techstat/>.

²⁰ *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>.

²¹ *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>.

²² For example, ISO/IEC 23894:2023 Information technology — Artificial intelligence — Guidance on risk management, <https://www.iso.org/standard/77304.html>.

Order. 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 the October 30, 2023 AI Executive Order. The practices in this section also do not supersede, modify, or direct an interpretation of existing requirements mandated by law or governmentwide policy, and agency responsible officials must coordinate to ensure that the performance 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 for one or more of the following purposes:
 - A. Evaluation of a potential vendor, commercial capability, or freely available AI capability that is not otherwise used in agency operations, solely for the purpose of making a procurement or acquisition decision;
 - B. Evaluation of a particular AI application because the AI provider is the target or potential target of a regulatory enforcement, law enforcement, or national security action;²³ and
 - C. Research and development.²⁴
- ii. **Extensions for Minimum Practices.** Until August 1, 2024, agencies may request from OMB an extension of limited and defined duration for a particular use of AI that cannot feasibly meet the minimum requirements in this section by that date. The request must be accompanied by a detailed justification for why the agency cannot achieve compliance for the use case in question and what practices the agency has in 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.
- 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 component²⁵ after making a written determination, based upon a system-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. Such waivers are applicable for the duration of the AI's use, but must be reassessed by the CAIO if there are significant changes to the conditions or context in which the AI is used. An agency CAIO may also revoke a previously issued waiver at

²³ Agencies are not required to follow these minimum practices when examining AI as the target or potential target of such an action, but they are required to follow these practices when *carrying out* an enforcement or national security action. For example, when evaluating an AI tool to determine whether it violates the law, agencies need not follow the minimum practices; if agencies were using that same tool to assess a different target, they would have to follow the minimum practices.

²⁴ AI research and development is not excluded if it is used in agency operations other than for the purposes of research and development, such as to make agency recommendations or decisions about real people.

²⁵ CAIOs may also grant waivers applicable to groups of closely related AI applications or components, provided that: (1) those systems have undergone a risk assessment that adequately considers the risks from each individual system; and (2) the systems are substantially identical in their risk profiles.

any time. Agencies must report to OMB within 30 days of granting such a waiver, detailing the scope, justifications, and supporting evidence.

iv. **Minimum Practices for Either Safety-Impacting or Rights-Impacting AI.**

Starting on August 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.** Impact assessments must document 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 an agency’s mission, for example to reduce costs, wait time for customers, or risk to human life, that can be measured after the AI is deployed to confirm or disprove the value of using AI.²⁶ Where quantification is not feasible, qualitative analysis should demonstrate an expected positive outcome, such as for improvements to customer experience or human interactions—and demonstrate that AI is a good fit to accomplish the relevant task.
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 reduce these risks. Agencies should document the stakeholders²⁷ that 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. If the agency cannot access such data after a reasonable effort to do so, it must obtain sufficient descriptive information from the AI or data provider to satisfy the

²⁶ 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.

²⁷ Stakeholders will vary by use case. 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.

reporting requirements in this paragraph. At a minimum, agencies must document:

- a. the provenance and quality of the data for its intended purpose;²⁸
- b. how the data is relevant to the task being automated and has a reasonable expectation of being useful for the AI's development, testing, and operation;
- c. whether the data contains sufficient breadth to address the range of real-world inputs the AI might encounter;
- d. whether the data comes from an adequately reliable source; and
- e. how errors from data entry, machine processing, or other sources are adequately measured and limited, to include errors from relying on AI-generated data as training data or model inputs.

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 customers that use the service who impact 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, to the extent practicable, that the AI will achieve its expected benefits while sufficiently mitigating risks associated with the AI, or else the agency should not use the AI. 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, both referenced in Section 5(c)(iv). Agencies must incorporate this independent evaluation into an applicable release or oversight process, or the Authorization to Operate process. The independent reviewing authority must not have been directly involved in the system's development.

²⁸ 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>, and 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.

Starting on August 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 and establish thresholds for periodic human review.** In addition to pre-deployment testing, agencies must institute ongoing procedures to monitor degradation to the AI's functionality and to detect changes in the AI's impact on rights or safety. Part of this monitoring process must include periodic human reviews to determine whether the existing implementation of the minimum practices in this section adequately mitigates any new risk. Such human review, including renewed testing for performance of the AI in a real-world context, must be conducted at least annually²⁹, and after significant modifications to the AI or to the conditions or context in which the AI is used. Reviews must include oversight and consideration by an appropriate internal agency authority not directly involved in the system's development or operation. 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 also monitor and defend the AI from AI-specific exploits,³⁰ particularly those that would adversely impact rights or safety.
- E. Mitigate emerging risks to rights and safety.** Upon identifying new or significantly altered risks to rights or safety through continuous 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 non-technical mitigations, such as greater human oversight. 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 is not practicable, agencies must stop using the affected AI as soon as is practicable.³¹
- F. 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 use case, product, or service being operated.

²⁹ For customer-facing services, agencies should consider customer feedback.

³⁰ For example, the AI-specific exploits outlined in the MITRE ATLAS framework. See <https://atlas.mitre.org/>.

³¹ 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.

- G. **Provide appropriate human consideration as part of decisions that pose a high risk to rights or safety.** Agencies should identify AI functionality that plays a role in decisions that pose a high risk to rights or safety and ensure that the AI functionality is not permitted to intervene directly in such situations without appropriate human consideration and accountability.
 - H. **Provide public notice and plain-language documentation through the AI use case inventory.** Agencies must ensure, to the extent consistent with applicable law and governmentwide guidance, including those 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 serves as adequately detailed and generally accessible documentation of the system's functionality that provides public notice of the AI to its users and the general public. Where practicable, agencies should include this documentation or link to it in contexts where people will interact with or be impacted by the AI. Where agencies' use cases are excluded from the public inventory requirements described in this guidance, 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.**
Starting on August 1, 2024, agencies must follow the above minimum practices for AI that is *either* safety-impacting *or* rights-impacting. In addition, agencies must also follow these minimum practices *before* initiating use of new or existing rights-impacting AI:
- A. **Take steps to ensure that the AI will advance equity, dignity, and fairness.**
This should include at least:
 - 1. *Proactively identifying and removing factors contributing to algorithmic discrimination or bias.* Agencies must assess whether their rights-impacting AI materially relies on information about a class protected by Federal nondiscrimination laws in a way that could result in algorithmic discrimination or bias against that protected class. Agencies should also assess whether proxies produce undue influence on their rights-impacting AI. In either case, if the AI's reliance on such information results in unlawful discrimination or harmful bias against protected classes, the agency must cease the use of the information before using the AI for decision-making.
 - 2. *Assessing and mitigating disparate impacts.* Agencies must test their AI to determine whether there are significant disparities in the AI's performance across demographic groups, including in the AI's real-world deployment,

and, consistent with applicable law, appropriately address disparities that have the potential to lead to discrimination, cause meaningful harm, or decrease equity, dignity, or fairness. If adequate mitigation of the disparity is not possible, then agencies should not use or integrate the AI tool.

3. *Using representative data.* Agencies should ensure that data used to develop, operate, and assess their AI is adequately representative of the communities who will be affected by the AI, and has been reviewed for improper bias based on the historical and societal context of the data.

B. Consult and incorporate feedback from affected groups. To the extent practicable and consistent with applicable law and governmentwide guidance, agencies must consult affected groups, including underserved communities, in the design, development, and use of the AI, and use such feedback to inform agency decision-making regarding the AI. In the event of negative feedback, agencies must consider not deploying the AI or removing the AI from use. Agencies are strongly encouraged to solicit feedback on an ongoing basis from affected groups, such as customers,³² Federal employee groups, and employees' union representatives, particularly after significant modifications to the AI or the conditions or context in which it is used. To carry out such consultations, agencies should take adequate steps to solicit input from the groups affected by the AI, which could include:³³

1. Direct user 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 open ended space for responses;
3. Post-transaction customer feedback collections;³⁴
4. Public hearings or meetings, such as a listening session; or
5. Any other transparent process that seeks public input, comments, or feedback from the affected groups in a meaningful, equitable, accessible, and effective manner.

Starting on August 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:

³² Customers can include individuals, businesses, or organizations that interact with an agency.

³³ Agencies are not required to conduct consultations in a format that would require OMB clearance under the Paperwork Reduction Act (44 U.S.C. § 3507), provided the steps the agency takes are adequate to solicit input from the groups affected by the AI.

³⁴ 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>.

- C. **Conduct ongoing monitoring and mitigation for AI-enabled discrimination.** As part of their ongoing monitoring requirement cited in Section 5(c)(iv)(D), agencies must also monitor rights-impacting AI to assess and mitigate AI-enabled discrimination against protected classes 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 affected AI functionality.
- D. **Notify negatively affected individuals.** Where practicable and consistent with applicable law and governmentwide guidance, agencies must notify individuals when AI meaningfully influences the outcome of decisions specifically concerning them, such as the denial of benefits.³⁵ Such notice should be timely and written in a manner that is consistent with the Plain Writing Act of 2010,³⁶ if applicable. 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 appropriate, requesting timely remediation for any related issues. Agencies are also strongly encouraged to provide explanations for such decisions and actions.³⁷
- E. **Maintain human consideration and remedy processes.** Agencies must provide timely human consideration and potential remedy to the use of the AI by 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. In developing appropriate remedies, agencies should follow OMB guidance on calculating administrative burden and the remedy process should not place unnecessary burden on the impacted individual.³⁸ When law or governmentwide guidance precludes disclosure of the use of AI or an opportunity for an individual appeal, agencies must create appropriate mechanisms for human oversight of rights-impacting AI.
- F. **Maintain options to opt-out where practicable.** Agencies must prominently provide and maintain a mechanism to conveniently opt out from AI functionality

³⁵ In some instances, such as an active law enforcement investigation, providing immediate notice may be inappropriate or impractical, and disclosure may be more appropriate at a later stage (i.e., prior to a defendant’s trial).

³⁶ Pub. L. No. 111-274 (codified at 5 U.S.C. § 301 note), <https://www.congress.gov/111/plaws/publ274/PLAW-111publ274.pdf>.

³⁷ Explanations might include, for example, how and why the AI-driven decision or action was taken. While exact explanations of AI decisions are often not technically feasible, agencies should 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.

³⁸ See OMB [M-22-10](#) and supporting document “[Strategies for Reducing Administrative Burden in Public Benefit and Service Programs](#).”

in favor of a human alternative where practicable and consistent with applicable law and governmentwide guidance. An opt-out mechanism must exist where the affected people have a reasonable expectation of an alternative or where lack of an alternative would meaningfully limit accessibility or create unwarranted harmful impacts.

d. Managing Risks in Federal Procurement of Artificial Intelligence

This section provides agencies with recommendations for responsible Federal procurement of 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 the AI Executive Order, OMB will also develop an initial means to ensure that AI contracts align with the guidance in this memorandum.

- i. **Aligning to National Values and Law.** Agencies should ensure that procured AI exhibits due respect for our Nation’s values, is consistent with the Constitution, and complies with all other applicable laws, regulations, and policies, including those addressing privacy, confidentiality, copyright, 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 of procured AI, such as through the use of model, data, and system cards;
 - B. regularly evaluating AI-performance claims made by Federal contractors, including in the particular environment where the agency expects to deploy the capability; and
 - C. considering contracting provisions that incentivize the continuous improvement of procured AI.
- 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 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 modifications 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 commercial AI offerings offered by the vendor without express permission from the agency.

- v. **Responsibly Procuring Generative AI.** Agencies are encouraged to include tailored risk management requirements in contracts for generative AI, and particularly for dual-use foundational models, including:
- A. requiring adequate testing and safeguards, including external AI red teaming, against risks from generative AI such as discriminatory, misleading, inflammatory, unsafe, or deceptive outputs;
 - B. 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
 - C. Agencies are encouraged to consider the relevant NIST standards, as appropriate, defined pursuant to Sections 4.1(a) and 10.1(d) of the AI Executive Order when imposing such requirements.

6. DEFINITIONS

The below definitions apply for the purposes of this memorandum.

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

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

Artificial Intelligence (AI): The term “artificial intelligence” has the meaning established in Section 238(g) of the John S. McCain National Defense Authorization Act for Fiscal Year 2019,³⁹ 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.
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.

³⁹ Pub. L. No. 115-232, § 238(g), <https://www.govinfo.gov/content/pkg/PLAW-115publ232/pdf/PLAW-115publ232.pdf>.

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 may assist in interpreting this definition:

1. This definition of AI encompasses, but is not limited to, the AI technical subfields of machine learning (including, but not limited to, 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, the technical complexity of a system (e.g., the number of parameters in a model, the type of model, or the amount of data used for training purposes) is not a relevant consideration for determining whether it constitutes AI.
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.

Artificial Intelligence Maturity: 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.

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

Automation Bias: 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.

CFO Act Agency: Refers to the agencies identified in 31 U.S.C. § 901(b).

Dual-Use Foundation Model: Has the meaning established in Section 3(k) of the AI Executive Order.

Equity: Has the meaning established in Section 10(a) of Executive Order 14091.⁴⁰

Federal Information: Has the meaning established in OMB Circular A-130.

Generative AI: Has the meaning established in Section 3(p) of the AI Executive Order.

⁴⁰ 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>.

Intelligence Community: Has the meaning established in 50 U.S.C. § 3003.

National Security System: Has the meaning established in 44 U.S.C. § 3552(b)(6).

Research and Development: As in OMB Circular No. A-11, *Preparation Submission, and Execution of the Budget* (2023), research and development is defined as creative and systematic work undertaken in order to increase the stock of knowledge—including knowledge of people, culture, and society—and to devise new applications using available knowledge.

Rights-Impacting AI:⁴¹ AI whose output serves as a basis for decision or action that has a legal, material, or similarly significant effect on an individual's or community'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, credit, employment, and other programs where civil rights and equal opportunity protections apply; or
3. Access to critical resources or services, including healthcare, financial services, social services, transportation, non-deceptive information about goods and services, and government benefits or privileges.

Risks from the Use of AI: 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;

⁴¹ Section 5(b) of this memorandum lists AI applications that are presumed to be rights-impacting.

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
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:⁴² AI that has the potential to meaningfully 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 21⁴³ and the infrastructure for voting and protecting the integrity of elections; or,
4. Strategic assets or resources, including high-value property, information marked as sensitive or classified by the Federal Government, and intellectual property.

Significant Modification: 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: Has the meaning established in Section 10(b) of Executive Order 14091.

⁴² Section 5(b) of this memorandum lists AI applications that are presumed to be safety-impacting.

⁴³ Presidential Policy Directive 21 (PPD-21), *Critical Infrastructure Security and Resilience*, <https://obamawhitehouse.archives.gov/the-press-office/2013/02/12/presidential-policy-directive-critical-infrastructure-security-and-resil>.

Appendix I: 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 Boards	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 (for DoD: submit to OMB metrics on use cases other than National Security Systems)	3(a)(iv), 3(a)(v)	Annually
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)	August 1, 2024 (with extensions possible)
Each Agency*	Report to OMB any agency-specific lists of AI purposes that are presumed to be rights-impacting or safety-impacting	5(b)	Annually
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 agencies in the Intelligence Community.