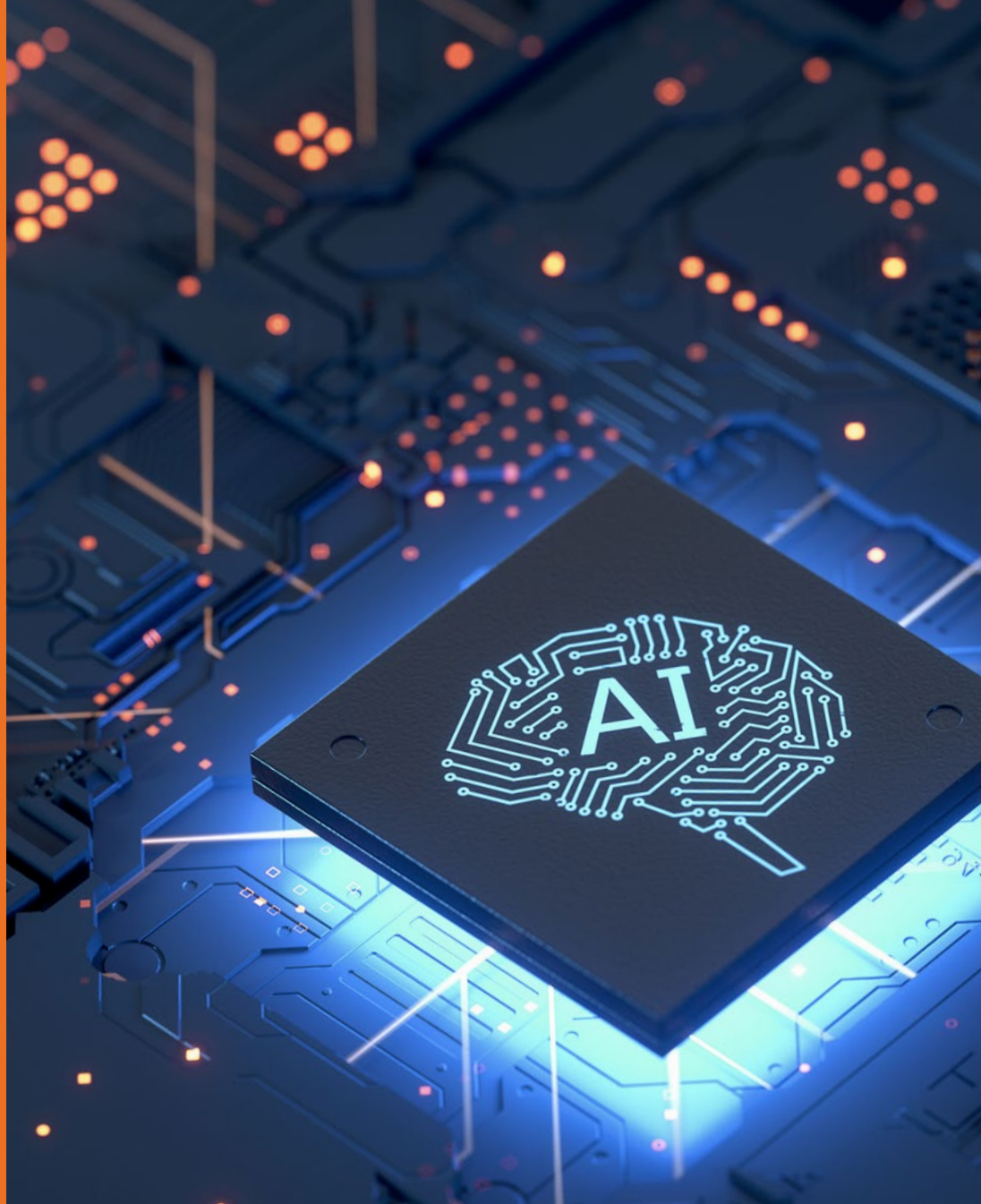




2023
2024

CISA ROADMAP

FOR

ARTIFICIAL INTELLIGENCE

PUBLICATION:
NOVEMBER 2023

CYBERSECURITY AND INFRASTRUCTURE SECURITY AGENCY



CONTENTS

INTRODUCTION1

VISION 2

CISA'S ROLE IN SECURING AI 2

FIVE LINES OF EFFORT 3

- **LINE OF EFFORT 1:** Responsibly Use AI to Support our Mission 5
- **LINE OF EFFORT 2:** Assure AI Systems7
- **LINE OF EFFORT 3:** Protect Critical Infrastructure From Malicious Use of AI . . 9
- **LINE OF EFFORT 4:** Collaborate with and Communicate on Key AI Efforts with the Interagency, International Partners, and the Public 11
- **LINE OF EFFORT 5:** Expand AI Expertise in our Workforce13

CONCLUSION14

KEY DEFINITIONS15

- Artificial Intelligence (AI)15
- AI Assurance15
- AI Security16
- Red Teaming16
- Adversarial Machine Learning16

APPENDIX Recent U.S. Efforts on AI Policy17

INTRODUCTION

As noted in the landmark Executive Order 14110, “Safe, Secure, And Trustworthy Development and Use of Artificial Intelligence (AI),” signed by the President on October 30, 2023, “AI must be safe and secure.” As the nation’s cyber defense agency and the national coordinator for critical infrastructure security and resilience, CISA will play a key role in addressing and managing risks at the nexus of AI, cybersecurity, and critical infrastructure.

This “2023–2024 CISA Roadmap for Artificial Intelligence” serves as a guide for CISA’s AI-related efforts, ensuring both internal coherence as well as alignment with the whole-of-government AI strategy. This roadmap incorporates key CISA-led actions as directed by Executive Order 14110, along with additional actions CISA is leading to promote AI security and support critical infrastructure owners and operators as they navigate the adoption of AI.

The roadmap includes CISA’s efforts to:

- Promote beneficial uses of AI to enhance cybersecurity capabilities and other aspects of CISA’s mission;
- Protect the nation’s AI systems from cybersecurity threats; and
- Deter malicious actors’ use of AI capabilities to threaten critical infrastructure.

The security challenges associated with AI parallel cybersecurity challenges associated with previous generations of software that manufacturers did not build to be secure by design, putting the burden of security on the customer. Although AI software systems might differ from traditional forms of software, fundamental security practices still apply. Thus, CISA’s AI roadmap builds on the agency’s cybersecurity and risk management programs. Critically, manufacturers of AI systems must follow secure by design principles: taking ownership of security outcomes for customers, leading product development with radical transparency and accountability, and making secure by design a top business priority. As the use of AI grows and becomes increasingly incorporated into critical systems, security must be a core requirement and integral to AI system development from the outset and throughout its lifecycle.

VISION

We envision a future in which AI systems advance our nation's cyber defense, where our critical infrastructure is resilient and protected from malicious use of AI, and where AI developers prioritize the security of their products as a core business requirement.

CISA'S ROLE IN SECURING AI

CISA's Strategic Plan 2023–2025 underpins CISA's adaptation to these technologies and each of CISA's four strategic goals are relevant to and impacted by AI:

GOAL 1 | CYBER DEFENSE. AI tools can help defend cyberspace against traditional threats, as well as emerging AI-driven threats. However, AI-based software systems are also software systems that require securing and necessitate cyber defense for AI.

GOAL 2 | RISK REDUCTION AND RESILIENCE. Critical infrastructure organizations increasingly use AI systems to maintain and improve resilience. CISA will guide and support responsible and risk-aware adoption of AI-based software systems that are secure by design.

GOAL 3 | OPERATIONAL COLLABORATION. As AI contributes to a rapidly changing threat landscape, CISA will communicate threat and risk information to the U.S. public, including critical infrastructure sectors. Furthermore, AI companies and AI use cases may be subject to targeted threats and may require specific services and protections in response.

GOAL 4 | AGENCY UNIFICATION. CISA will responsibly integrate AI software systems across the agency, as well as recruit and develop a workforce capable of optimally harnessing AI software systems to carry out CISA's mission.

FIVE LINES OF EFFORT

This roadmap represents our work to unify and accelerate CISA's AI efforts along five lines of effort (LOE):

LINE OF EFFORT 1:

Responsibly use AI to support our mission. CISA will use AI-enabled software tools to strengthen cyber defense and support our critical infrastructure mission. CISA's adoption of AI will ensure responsible, ethical, and safe use—consistent with the Constitution and all applicable laws and policies, including those addressing federal procurement, privacy, civil rights, and civil liberties.

LINE OF EFFORT 2:

Assure AI systems. CISA will assess and assist secure by design, AI-based software adoption across a diverse array of stakeholders, including federal civilian government agencies; private sector companies; and state, local, tribal, and territorial (SLTT) governments through the development of best practices and guidance for [secure and resilient AI software development and implementation](#).

LINE OF EFFORT 3:

Protect critical infrastructure from malicious use of AI. CISA will assess and recommend mitigation of AI threats facing our nation's critical infrastructure in partnership with other government agencies and industry partners that develop, test, and evaluate AI tools.

LINE OF EFFORT 4:

Collaborate with and communicate on key AI efforts with the interagency, international partners and the public. CISA will contribute to DHS-led and interagency processes on AI-enabled software. This LOE includes developing policy approaches for the U.S. government's overall national strategy on AI and supporting a whole-of-DHS approach on AI-based-software policy issues. This LOE also includes coordinating with international partners to advance global AI security best practices and principles.

LINE OF EFFORT 5:

Expand AI expertise in our workforce. CISA will continue to educate our workforce on AI software systems and techniques, and the agency will continue to actively recruit interns, fellows, and future employees with AI expertise. CISA will ensure that internal training reflects—and new recruits understand—the legal, ethical, and policy aspects of AI-based software systems in addition to the technical aspects.



Figure 1. CISA's roadmap for artificial intelligence lines of effort

This roadmap provides objectives for each line of effort that detail how CISA will accomplish these goals and measure our success. We also include representative outcomes and a notional measurement approach for each line of effort. We are developing more specific measures of effectiveness, which will be defined in our annual operating plans. Of note, identifying appropriate measures of effectiveness and vice measurements of performance is challenging and will require an ongoing effort—with continuous refinements as needed—throughout the life of the plan.

LINE OF EFFORT 1

RESPONSIBLY USE AI TO SUPPORT OUR MISSION



CISA will use AI-enabled software tools to strengthen cyber defense and support our critical infrastructure mission. CISA's adoption of AI will ensure responsible, ethical, and safe use—consistent with the Constitution and all applicable laws and policies, including those addressing federal procurement, privacy, civil rights, and civil liberties.

REPRESENTATIVE OUTCOMES

1 | CISA assesses our cybersecurity programs for potential uses of AI and provides the resources, requirements, and oversight to incorporate AI where appropriate.

2 | Through the responsible use of AI tools, CISA network defenders proactively mitigate threats to critical networks before damaging intrusions occur.

MEASUREMENT APPROACH

Increased responsible uses of AI software tools across CISA workflows.

OBJECTIVE 1.1 | Establish governance and oversight processes for CISA's use of AI.

CISA will establish robust AI governance processes to coordinate actions across the agency. This will include developing ethical and safety oversight processes as well as legal, procurement, privacy, civil rights, and civil liberties considerations. Responsible use will be central to our application of AI.

To promote responsible AI use, CISA will:

- Create our own [NIST AI Risk Management Framework \(RMF\)](#) profile to help develop and implement security and privacy controls for AI;
- Implement a programmatic structure for AI adoption within cyber defense missions;
- Review active AI use cases;
- Develop workplace guidance for generative technologies; and,
- Address AI data requirements and uses.

OBJECTIVE 1.2 | Collect, review, and prioritize AI use cases to support CISA missions.

CISA will create an agency AI Use Case Inventory to collect, review, and prioritize AI use cases supporting our missions. This inventory will encompass improvements to existing IT systems, collaboration tools, workflows, critical infrastructure defense programs, and proposed data collections for training AI models.



OBJECTIVE 1.3 | Develop an adoption strategy for the next generation of AI-enabled technologies.

To stay ahead of the adoption curve while ensuring privacy and civil rights protections, CISA will closely coordinate AI-related research and development efforts to target gaps in mission needs. These initiatives include the identification of baseline responsible practices for AI to protect safety and rights, the creation of a safe and secure AI testbed, and the development of technical requirements for cybersecurity use cases.

OBJECTIVE 1.4 | Incorporate cyber defense, incident management, and redress procedures into AI systems and processes.

CISA will establish incident response capabilities for AI usage, including remedy and redress procedures when necessary. In addition, CISA will adopt an approach for continuous evaluation of AI models while reviewing IT security practices to securely integrate AI technology.

OBJECTIVE 1.5 | Examine holistic approaches to limiting bias in AI use at CISA.

CISA will explore holistic approaches to limit bias in AI use, identifying potential bias points in the development, testing, implementation, and maintenance processes in order to build in fairness. Beyond exploring bias and mitigation strategies, CISA will develop a quality assessment for training data and public notice of our AI Use Case Inventory.

OBJECTIVE 1.6 | Responsibly and securely deploy AI systems to support CISA's cybersecurity mission.

Responsible and secure AI deployment aligns with CISA's core cybersecurity mission. To help ensure this, CISA will explore the identification, testing, evaluation, and deployment of AI capabilities for cyber defense, including detection of vulnerabilities in critical U.S. government software, systems, and networks, and we will document the lessons learned.

LINE OF EFFORT 2

ASSURE AI SYSTEMS



CISA will assess and assist secure by design AI-based software adoption across a diverse array of stakeholders, including federal civilian government agencies; private sector companies; and state, local, tribal, and territorial (SLTT) governments through the development of best practices and guidance for secure and resilient AI software development and implementation.

REPRESENTATIVE OUTCOMES

1 | CISA identifies cybersecurity risks and security resilience challenges as early as possible during AI adoption to mitigate threats to critical infrastructure.

2 | CISA adapts existing security guidance and service offerings to AI software systems, including best practices for red teaming AI systems and for making AI software that is secure by design.

3 | Stakeholders understand how AI-specific vulnerabilities fit into the existing coordinated [vulnerability disclosure](#) process.

MEASUREMENT APPROACH

Increased adherence to CISA risk guidance and best practices for AI software deployment, including guidance on red teaming and vulnerability management.

OBJECTIVE 2.1 | Assess cybersecurity risks of AI adoption in critical infrastructure sectors.

CISA will assess potential risks related to the use of AI in critical infrastructure sectors, including ways in which deploying AI may make critical infrastructure systems more vulnerable to critical failures, physical attacks, and cyberattacks. CISA will then consider ways to mitigate these vulnerabilities. Additionally, CISA will incorporate the [NIST AI Risk Management Framework \(AI RMF 1.0\)](#), as well as other appropriate security guidance, into relevant safety and security guidelines and best practices for use by critical infrastructure owners and operators.

OBJECTIVE 2.2 | Engage critical infrastructure stakeholders to determine security and resilience challenges of AI adoption.

CISA will engage with critical infrastructure stakeholders to assess and address the use of AI across critical infrastructure sectors.

OBJECTIVE 2.3 | Capture the breadth of AI systems used across the federal enterprise.

CISA will evaluate [Software Bill of Materials \(SBOM\)](#) toolchains, including SBOM format standards and automated SBOM collection and translation software, to confirm coverage of AI software.



OBJECTIVE 2.4 | Develop best practices and guidance for acquisition, development, and operation of secure AI systems.

CISA will develop best practices and guidance for the acquisition, development, and operation of secure AI systems. We will also provide guidance for the secure use of AI technologies and will integrate this guidance into the [Cybersecurity Performance Goals](#) pertaining to AI and related systems.

OBJECTIVE 2.5 | Drive adoption of strong vulnerability management practices for AI systems.

CISA will develop tools and techniques to harden and test AI systems, as well as incorporate appropriate outputs of adversarial ML processes and AI system vulnerabilities into the [National Vulnerability Database](#). This includes conducting an operational test of an AI vulnerability in the [Coordinated Vulnerability Disclosure \(CVD\)](#) process, as well as writing strategic guidance for security testing and red teaming AI systems and software, particularly [Open Source Software](#).

OBJECTIVE 2.6 | Incorporate AI systems into Secure by Design initiative.

CISA champions a secure by design approach to developing and manufacturing technology products, ensuring manufacturers design products with security in mind at the onset, so consumers receive products that are secure right out of the box. To encourage a secure by design approach to AI software and products, CISA will integrate AI security into the [Secure by Design](#) program and will develop a research pipeline to continually understand and project ways to support AI systems security.

LINE OF EFFORT 3

PROTECT CRITICAL INFRASTRUCTURE FROM MALICIOUS USE OF AI



CISA will assess and recommend mitigation of AI threats against our nation's critical infrastructure in partnership with other government agencies and industry partners that develop, test, and evaluate AI tools.

REPRESENTATIVE OUTCOMES

- 1 | Through engagement with stakeholders, including tabletop exercises focused on AI-enhanced attacks, CISA protects AI systems from adversarial manipulation or abuse.
- 2 | CISA supports the advancement of AI risk management practices across the critical infrastructure community through the publication and dissemination of decision support materials, such as a risk management guide for AI risks to critical infrastructure.

MEASUREMENT APPROACH:

Number of publications and engagements that support shared awareness of emerging AI-related risks and advances in AI risk management practices.

OBJECTIVE 3.1 | Regularly engage industry stakeholder partners that are developing AI tools to assess and address security concerns to critical infrastructure and evaluate methods for educating partners and stakeholders.

CISA will build on existing structures to advance industry collaboration and coordination around AI security. The Information Technology Sector Coordinating Council's AI Working Group, which was established in March 2023, will continue to provide advice on AI security challenges and feedback on agency AI initiatives.

CISA will also stand up an operational effort in the [Joint Cyber Defense Collaborative \(JCDC\)](#), JCDC.AI, to catalyze focused collaboration around threats, vulnerabilities, and mitigations affecting AI systems. The JCDC effort will also explore potential operational planning efforts that bring together AI providers and critical infrastructure operators to address specific risks.



OBJECTIVE 3.2 | Use CISA partnerships and working groups to share information on AI-driven threats.

CISA will use agency partnerships and working groups, including JCDC.AI, to share information on AI-driven threats. The agency will engage industry, federal, and international partners to understand emerging threats and share them with the broader community.

OBJECTIVE 3.3 | Assess AI risks to critical infrastructure.

Each critical infrastructure sector has a unique set of needs and capabilities. As adversaries adopt AI-enabled software systems and as AI expands the cyber threat landscape, CISA will publish materials to raise awareness of emerging risks. CISA will also evaluate risk management approaches and methodologies to determine the appropriate analytical framework for the assessment and treatment of AI risks and will identify necessary enhancements.

LINE OF EFFORT 4



COLLABORATE WITH AND COMMUNICATE ON KEY AI EFFORTS WITH THE INTERAGENCY, INTERNATIONAL PARTNERS, AND THE PUBLIC

CISA will contribute to DHS-led and interagency processes on AI-enabled software. This LOE includes developing policy approaches for the U.S. government's overall national strategy on AI and supporting a whole-of-DHS approach on AI-based-software policy issues. This LOE also includes coordinating with international partners to advance global AI best practices and principles.

REPRESENTATIVE OUTCOMES

1 | CISA stakeholders are aligned around clear guidance for AI security.

MEASUREMENT APPROACH

Proportion of AI-focused guidance and policy documents developed in collaboration with U.S. interagency and international partners.

OBJECTIVE 4.1 | Support the development of a whole-of-DHS approach on AI policy issues.

CISA will support the development of a whole-of-DHS approach to AI policy issues. As CISA develops our agency-specific AI efforts, we will closely coordinate with DHS entities, including the DHS AI Task Force.

OBJECTIVE 4.2 | Participate in interagency policy meetings and interagency working groups on AI.

CISA will attend interagency meetings to foster coherent and collaborative approaches to federal government AI policy.

OBJECTIVE 4.3 | Develop CISA policy positions that take a strategic, national level perspective for AI policy documents, such as memoranda and other products.

CISA will develop policy positions that take a strategic, national level perspective for AI policy documents and ensure alignment of CISA strategies, priorities, and policies with interagency doctrine. CISA will drive policy decisions to support critical infrastructure equities and integrate national strategic level perspectives in key AI policy documents. Additionally, to increase public awareness about AI assurance, CISA will develop AI assurance publications.



OBJECTIVE 4.4 | Ensure CISA strategy, priorities, and policy framework align with interagency policies and strategy.

CISA will work across the interagency to ensure CISA policies and strategies align with the whole-of-government approach to AI.

OBJECTIVE 4.5 | Engage with international partners surrounding global AI security.

CISA will co-develop and co-seal guidance for AI security with other federal agencies and international partners. CISA will engage with international partners surrounding global AI security and encourage the adoption of international best practices for secure AI.

LINE OF EFFORT 5

EXPAND AI EXPERTISE IN OUR WORKFORCE



CISA will continue to educate our workforce on AI software systems and techniques, and the agency will continue to actively recruit interns, fellows, and future employees with AI expertise. CISA will ensure that internal training reflects—and new recruits understand—the legal, ethical, and policy aspects of AI-based software systems in addition to the technical aspects.

REPRESENTATIVE OUTCOMES

1 | CISA hires, trains, and retains a workforce with AI expertise.

MEASUREMENT APPROACH

Increased AI expertise in the CISA workforce.

OBJECTIVE 5.1 | Connect and amplify AI expertise that already exists in CISA's workforce.

As the nation's cyber defense agency, the CISA team includes a strong workforce of cybersecurity experts. The agency will identify and leverage existing AI expertise across CISA. We will also develop an AI community of practice for engagement across the agency, as well as maintain key points of contact from each division leading AI activities, positioning the agency for a collaborative and cohesive approach to expanding our AI capabilities.

OBJECTIVE 5.2 | Recruit interns, fellows, and staff with AI expertise.

CISA will recruit interns, fellows, and staff with AI expertise. CISA will use a variety of pathways, including the Cyber Talent Management System (CTMS), for recruiting, developing, and maintaining our AI workforce.

OBJECTIVE 5.3 | Educate CISA's workforce on AI.

CISA will provide training and education opportunities for employees on an ongoing basis as part of our plan to help our workforce have the knowledge and skills to engage, innovate, and apply appropriately the current and emerging capabilities afforded by AI.

OBJECTIVE 5.4 | Ensure internal training not only reflects technical expertise, but also incorporates legal, ethical, and policy considerations of AI implementation across all aspects of CISA's work.

CISA will provide access to training that includes objectives on legal, ethical, and policy aspects of implementing AI.

CONCLUSION

A whole-of-government approach is needed to fully harness the benefits and mitigate the risks of AI. Through the initiatives outlined in this roadmap, CISA strives toward our vision of a nation in which AI systems advance our nation's cyber defense, where our critical infrastructure is resilient and protected from malicious use of AI, and where AI developers prioritize the security of their products as a core business requirement.

KEY DEFINITIONS

ARTIFICIAL INTELLIGENCE (AI)

Within this document, “Artificial intelligence” (AI) has the meaning¹ set forth in the *National Artificial Intelligence Initiative Act of 2020* (enacted as Division E of the William M. (Mac) Thornberry National Defense Authorization Act for Fiscal Year 2021 (Public Law 116-283), Section 5002(3):

A machine-based system that can, for a given set of human-defined objectives, make predictions, recommendations or decisions influencing real or virtual environments. Artificial intelligence systems use machine and human-based inputs to:

- (A) perceive real and virtual environments;
- (B) abstract such perceptions into models through analysis in an automated manner;
- and,
- (C) use model inference to formulate options for information or action.

AI encompasses machine learning (ML), which, according to Executive Order 14110 is “a set of techniques that can be used to train AI algorithms to improve performance on a task based on data.”

AI ASSURANCE

Many terms are shared between the AI and information security communities, but the same term can carry different and incompatible meanings. Because this roadmap is relevant to both communities, this document incorporates both sets of meanings. As an example, both communities have developed special uses of “assurance” independently since at least the 1980s:

AI ASSURANCE: “A process that is applied at all stages of the AI engineering lifecycle ensuring that any intelligent system is producing outcomes that are valid, verified, data-driven, trustworthy and explainable to a layman, ethical in the context of its deployment, unbiased in its learning, and fair to its users.”²

SECURITY ASSURANCE: “Measure of confidence that the security features, practices, procedures, and architecture of an information system accurately mediate and enforce the security policy.”³

When this document simply says “assurance” or uses another shared term without distinguishing the origin, this document incorporates **both** communities’ meanings.

¹ There are other statutory definitions for artificial intelligence. The “AI in Government Act of 2020” (P.L. 116-260, Division U, Title I, codified at 40 U.S.C. § 11301, note), listed earlier in this document, uses the definition from § 238(g) of the John S. McCain National Defense Authorization Act for Fiscal Year 2019, (P.L. 116-232, codified at 10 U.S.C. § 2358 note).

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.
- (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.

² Batarseh FA, Freeman L, Huang CH. A survey on artificial intelligence assurance. *Journal of Big Data*. 2021 Apr 26;8(1):60. [A survey on artificial intelligence assurance | Journal of Big Data \(springer.com\)](#)

AI SECURITY

AI Security is a term encompassing several different categories of cybersecurity, including the three key categories addressed in this roadmap:

1. Applications of AI for cybersecurity:

CISA actively leverages AI tools for threat detection, prevention, and vulnerability assessments.

2. Cybersecurity of AI-enabled systems:

CISA is applying traditional cybersecurity principles and practices to protect and secure AI-enabled systems. In addition to being well-positioned to leverage our existing expertise, CISA is advancing AI-enabled systems security through efforts to promote [secure by design](#) best practices for AI-enabled software systems.

3. Threats from Adversarial Use of AI:

CISA, in collaboration with other parts of DHS, will focus on research, development, and acquisition of tools to improve the resilience of federal civilian executive branch (FCEB) agencies and critical infrastructure and to protect these agencies and organizations from malicious uses of AI.

RED TEAMING

As defined in Executive Order 14110, AI red teaming is “a structured testing effort to find flaws and vulnerabilities in an AI system, often in a controlled environment and in collaboration with developers of AI systems. AI red-teaming is most often

performed by dedicated ‘red teams’ that adopt adversarial methods to identify flaws, vulnerabilities, or logic errors, such as harmful or discriminatory outputs from an AI system, unforeseen or undesirable system behaviors, limitations, or potential risks associated with the misuse of the system.”

The goals of security testing and red teaming (whether on AI systems or not) are to identify vulnerabilities in a system and better manage the associated security posture. During a cybersecurity [red team assessment](#), a red team attempts to gain access to an organization’s enterprise network and trigger a security response from the organization’s people, processes, or technology.

ADVERSARIAL MACHINE LEARNING

Malicious cyber actors target vulnerabilities throughout the AI supply chain to cause certain behavior, unintended by the system owner or operator, in machine learning (ML) systems—referred to as adversarial machine learning. For example, malicious actors may manipulate training data, affect the performance of the ML model’s classification and regression, or exfiltrate sensitive ML model information. For more information on adversarial machine learning, including information on types of activity, see National Institute of Standards and Technology (NIST), Technical Series Publication [Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations](#) and MITRE ATLAS [Adversarial Machine Learning 101](#).

³ NIST SP 800-39, [Managing Information Security Risk Organization, Mission, and Information System View](#). This definition is also very similar to the international IETF definition in [RFC 4949](#).

⁴ CISA Cybersecurity Advisory: [CISA Red Team Shares Key Findings to Improve Monitoring and Hardening of Networks](#)

⁵ Poison Training Data. MITRE ATLAS™ (Adversarial Threat Landscape for Artificial-Intelligence Systems) Techniques. <https://atlas.mitre.org/techniques/AML.T0020/>.

APPENDIX

RECENT U.S. EFFORTS ON AI POLICY

Recent actions taken by the U.S. government's executive and legislative branches related to AI-based software systems reflect the need to marshal a national effort to defend critical infrastructure and government networks and assets, work with partners across government and industry, and expand existing services and programs for federal civilian agencies and critical infrastructure owners and operators. The following recent efforts guide CISA's actions in this plan:

Executive Order 14110 “Safe, Secure, And Trustworthy Development and Use of Artificial Intelligence (AI).” (October 2023) This EO focuses on ensuring that AI is safe and secure. This will require robust, reliable, repeatable, and standardized evaluations of AI systems, as well as policies, institutions, and mechanisms to test, understand, and mitigate risks from these systems before they are put to use.

Voluntary Commitments from Leading Artificial Intelligence Companies to Manage the Risks Posed by AI. (Updated September 2023) The Biden-Harris administration has secured voluntary commitments from leading AI companies to help move toward safe, secure, and transparent development of the AI technology. These commitments include ensuring products are safe before introducing them to the public, building systems that put security first, and earning the public's trust.

DHS Policy Statement 139-06 Acquisition and Use of Artificial Intelligence and Machine Learning by DHS Components. (August 2023) This policy statement provides that DHS will acquire and use AI only in a manner that is consistent with the Constitution and all other applicable laws and policies.

New National Science Foundation Funding. (May 2023) This dedicated \$140 million will launch seven new National AI Research Institutes to promote responsible innovation, bolster America's AI research and development (R&D) infrastructure and support the development of a diverse AI workforce.

AI Risk Management Framework (RMF). (January 2023) In collaboration with the private and public sectors, the National Institute of Standards and Technology (NIST) developed this framework to better manage risks—to individuals, organizations, and society—that are uniquely associated with AI. The NIST AI RMF, intended for voluntary use, aims to improve the ability to incorporate trustworthiness considerations into the design, development, use, and evaluation of AI products, services, and systems.

Blueprint for an AI Bill of Rights. (October 2022) This framework is a set of five principles—identified by the White House Office of Science and Technology Policy—that should guide the design, use, and deployment of automated systems to protect the American public in the age of AI.

2021 Final Report of the National Security Commission on Artificial Intelligence. (March 2021) This report presented an integrated national strategy to reorganize the government, reorient the nation, and rally our closest allies and partners to defend and compete in the coming era of AI-accelerated competition and conflict.

National Artificial Intelligence Initiative (NAII) Act of 2020 (Division E of the National Defense Authorization Act for Fiscal Year 2021). (January 2021) Among other things, this act established direction and authority to coordinate AI research, development, and demonstration activities among civilian agencies, the Department of Defense, and the intelligence community to ensure that each informs the work of the others.

AI in Government Act of 2020 (Title I of Division U of the Consolidated Appropriations Act, 2021). (December 2020) This act created the AI Center of Excellence within the General Services Administration and directed the Office of Management and Budget (OMB) to issue a memorandum informing federal agencies of policies for acquisition and application of AI and identifying best practices for mitigating risks.

Department of Homeland Security 2020 Artificial Intelligence Strategy. (December 2020) This strategy set out to enhance DHS’s capability to safeguard the American people, our homeland, and our values through the responsible integration of AI into DHS’s activities and the mitigation of new risks posed by AI.

EO 13960: Promoting the Use of Trustworthy AI in the Federal Government. (December 2020) This executive order required federal agencies to inventory their AI use cases and share their inventories with other government agencies and the public.

EO 13859: Maintaining American Leadership in AI. (February 2019) This executive order established federal principles and strategies to strengthen the nation’s capabilities in AI to promote scientific discovery, economic competitiveness, and national security.

2023–2024
CISA ROADMAP
— FOR —
ARTIFICIAL INTELLIGENCE

