



Securing the Agentic Enterprise: An Identity Security Approach

An executive view for leaders on adopting AI agents safely.

AI agents are moving into government work quickly. They can read and write data, use applications, and act across your systems at machine speed, with no person at the keyboard. Used well, they are a powerful force multiplier for the mission. Left ungoverned, each one is a new account with broad access and no owner, and a new way for things to go wrong. The reassuring part: securing AI is not a new discipline. It is identity, applied to machines. This guide explains what the risk is, what federal mandates already require, and the practical steps to adopt AI with confidence.

What this paper covers and what it does not.

This paper is about the identity security of AI agents: giving each agent a verifiable identity, scoping what it can access, carrying that identity safely across the systems it touches, and revoking it quickly when something goes wrong. Identity decides who and what an agent can act as, and it is where an agency has the most immediate, standards-based control today. It does not cover the other parts of using AI responsibly, such as the accuracy of model outputs, bias and fairness, content safety and the prevention of prompt injection, data privacy, or broad AI risk-management frameworks. Those matter and are addressed elsewhere. Where they meet identity – for example, an agent manipulated through its prompt or behaving unpredictably – strong identity and access controls limit the damage and preserve accountability, but preventing the manipulation itself is out of scope here.

[Our companion technical paper](#) names the specific standards and shows your architects and security engineers how to put them in place.

Why leaders must address this now.

An AI agent is a new kind of user on your network, a non-human one. It signs in on its own, uses tools and data, and can even create other agents to help it work. This is not unfamiliar territory: agencies have managed non-human accounts, the service accounts and automated identities behind everyday systems, for years. What is new is the speed and scale. Software and machine identities already outnumber human users¹ by more than 100:1 in a typical enterprise, and AI agents are the fastest-growing group. They are often created quickly, sometimes in large numbers,

and are easy to lose track of. Deciding how to govern them is a leadership and risk decision, not only an IT task, because it determines whether your agency can adopt AI safely or inherits a fast-growing blind spot.

What can go wrong.

Six problems recur when machine and agent identities are not governed. Each carries a real cost:

What can go wrong	What it means for you
Passwords that never expire	Machines often hold keys or passwords that sit unchanged for months. If one leaks, an attacker has a way in long after the agent is gone.
Too much access	Agents are often given broad permissions to make them useful, so a single compromised agent can reach far more than its job required.
Acting for the wrong person	A low-privilege user, or a manipulated prompt, can push a powerful agent to do something it should not, and the logs cannot tell who really acted.
Stolen access passed along	As agents hand work to other services, the access they pass can be stolen and reused to impersonate them.
No off switch	When access is granted once and never re-checked, a compromised agent keeps working until its access happens to expire.
No accountability	Shared or anonymous machine accounts mean actions cannot be traced to an owner, so audits and incident response break down.

What the mandates already require.

Securing AI is not a separate initiative. It is part of the Zero Trust direction federal agencies are already following. Three touchstones make this explicit:

Federal Zero Trust Strategy (OMB M-22-09). Access should be granted based on strong identity, least-privilege authorization, and continuous evaluation of access requests and active sessions.

Identity Lifecycle (OMB M-19-17). Agencies must manage the lifecycle of human and non-person identities, including mechanisms to bind, update, revoke, and destroy credentials for devices and automated technologies.

Zero Trust Architecture (NIST SP 800-207). Access decisions should be made on a least-privilege, per-request basis, with trust continually evaluated for both person and non-person entities.

Bottom line for leaders: AI agents are identities, and the zero trust rules you already follow apply to them directly.

Five ways to secure your AI.

Whatever platforms you run, securing AI agents comes down to five plain principles:

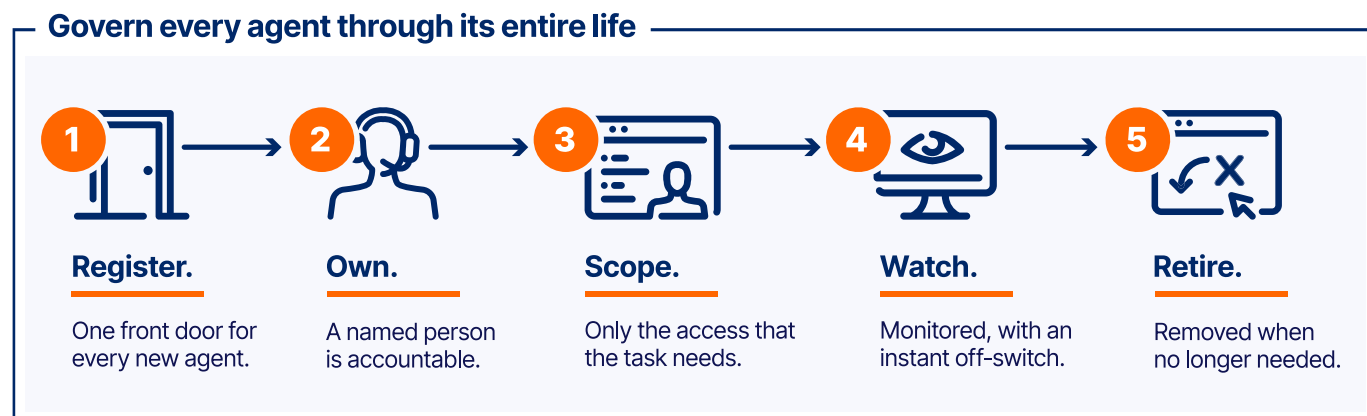
Give every agent an identity and an owner. No anonymous, shared, or abandoned agents. Every agent has a unique identity and a named person accountable for what it does.

No stored passwords for machines. Replace static keys and passwords with short-lived credentials that expire on their own, so there is nothing lasting to steal.

Least privilege, checked every time. Give each agent only the access its task needs, and verify that access every time it acts, not just at first sign-in.

Continuous oversight and an instant off switch. Watch what agents do, and keep the ability to shut a compromised agent off in seconds rather than waiting for its access to time out.

Govern the whole life of the agent. Register it, review it on a schedule, and retire it when it is no longer needed, so agents do not pile up unwatched.



Build on open standards and the platforms you already run.

Two practical choices make this durable. First, use open standards. The technology industry has agreed on common, open standards for machine and agent identity. Adopting them means your security controls work across different vendors, you avoid being locked into one product, and you are ready as the technology keeps changing. Second, build on what you already have. The identity platforms most agencies already run, such as Microsoft Entra and Okta, are adding these capabilities now. Adopt each for what it does well today, and plan for the parts that are still maturing. Your engineering teams do not need to start from scratch, and they should not build a separate, one-off system for AI.

A practical path.

You do not have to solve everything at once. The work moves through three stages, each building on the last. You can begin the first in a matter of weeks, and each stage delivers value on its own.

Stage 1: Get your arms around it. → First 90 days

The goal is simple: Know what you have, and make it someone's job.

Put one team in charge. Make AI-agent identity a named responsibility of your enterprise identity team, so it is owned rather than left to each project to work out alone.

Take inventory. Find the machine and agent identities already in your environment, and record what each one is, what it does, and who owns it.

Clean house. Assign an owner to anything that lacks one, and retire the accounts no one owns or uses. Removing these is the fastest risk reduction available to you.

Write down the ground rules. Agree what every agent must have before it goes live: a unique identity, a named owner, only the access it needs, no stored passwords, and a way to switch it off.

By the end of this stage, you have one list of every agent, an owner for each, and a clear definition of what governed means.

Stage 2: Build one front door. → Next 3-6 months

Now make the governed path the only path.

Stand up the front door. Create one place where every new agent is registered, given an owner, granted only the access its task needs, and placed under monitoring before it goes live.

Make it the only way in. Set the expectation, as a leadership decision, that agents reach production through the front door and not around it.

Use what you already own. Where your identity platform offers agent governance built in, use it rather than standing up a separate system to maintain.

Retire stored passwords. Move agents off static keys and passwords onto short-lived credentials that expire on their own, so there is nothing lasting to steal.

By the end of this stage, every new agent is governed from day one, holds no lasting secret, and traces back to an accountable owner.

Stage 3: Keep it healthy. → Ongoing

Turn governance into something that runs by default, not something you rediscover at audit time.

Review on a schedule. Have owners confirm, on a regular cadence, that each agent is still needed and still has the right access, and automate the reminders and removals so nothing lingers.

Watch, and be ready to act. Monitor what agents do, and keep the ability to switch off a compromised agent quickly.

Report what matters. Give leadership a few plain measures: how many agents are governed, how fast you can shut a bad one off, and how many unowned accounts remain.

Improve as you go. Feed new agents and new lessons back into the standard so the program keeps pace with adoption.

The result is a durable, auditable program in which agents are governed continuously, and you can prove it.

What good looks like.

When this is in place, leaders can answer the questions that matter. You can see every agent and who owns it. Agents carry only the access they need and can be switched off instantly if they misbehave. You can adopt AI fast enough to deliver on the mission without opening a new class of standing risk. And you can show auditors that your AI is governed under the same Zero Trust rules as everything else.

How Easy Dynamics helps.

Easy Dynamics designs, builds, and secures the identity systems that protect federal missions. We help agencies put

this into practice: Standing up agent governance on the platforms you already run, cleaning up the machine identities you already have, and building the governed front door that lets you adopt AI with confidence. We meet you where you are, from strategy through implementation and steady-state operations.

To discuss securing AI at your agency, contact the Easy Dynamics Digital Identity and Cybersecurity practice at info@easydynamics.com.

¹ Platform capabilities and standards status reflect information available as of mid-2026 and continue to evolve. Identity chaining, transaction tokens, and SPIFFE client authentication are advancing through the IETF; Shared Signals and CAEP are finalized OpenID Foundation specifications. Microsoft Agent 365 and Entra Agent ID reached general availability in May 2026. Broad product support for several of these standards, continuous authorization in particular, will be limited for the foreseeable future; treat this as a multi-year target and re-assess as support matures.

For the engineering detail behind this guide, see the companion technical paper, [A Standards-Centric Identity Architecture for Federal AI Agents](#). Platform capabilities continue to evolve; this guide reflects the landscape as of 2026.