



A GOVERNMENT TECHNOLOGY THOUGHT LEADERSHIP PAPER



Securing AI Agents in Government

SPONSORED BY



Agentic AI could represent the next frontier of government transformation.

According to a March 2026 report from the National Association of State Chief Information Officers (NASCIO), at least eight states have agentic AI projects in production.¹ Other state and local agencies are looking into or developing similar initiatives.

Although AI agents that complete processes by themselves can significantly boost government efficiency and effectiveness, adoption and experimentation should not outpace governance.

AI agents can spawn their own additional bots and connect across apps and data systems — often outside existing security controls. In addition to managing their own agents, agencies must fend off agents deployed by malicious actors.

Governments can contain these risks by treating AI agents as a new identity type governed with the same rigor applied to human users. Only through this approach can agencies securely benefit from agentic AI.

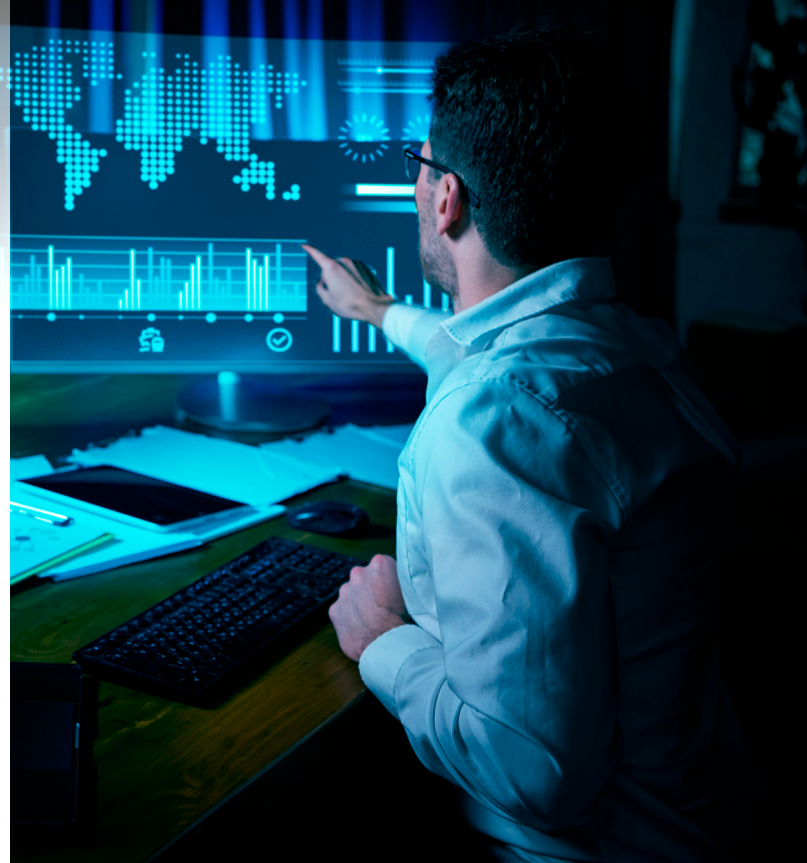
Emerging Challenges

Attractive use cases for agentic AI highlight the need for new machine identity controls.

A health and human service agency, for instance, might use AI agents to accelerate benefits approvals and reduce errors. “The agent connects to tax, health and benefits databases and touches personal identifiable information across all of them,” says Daniel Watts, product marketing manager for state and local government and education with Okta. “We need to know who authorized it and what credentials it’s using. Government must document everything.”

The following challenges call for comprehensive governance:

- ✓ AI agents can proliferate instantly and quietly at massive scale.
- ✓ Agentic AI risks compound if agencies lack visibility into who owns each agent, what data agents can access and what actions agents can take.
- ✓ Agentic initiatives can’t get off the ground without a clear strategy for managing agent identities.



3 Questions to Ask

1. Where are my agents?

IT leaders have spent decades building identity security measures around their workforce: directories, structured onboarding and offboarding, access reviews and so on. That scaffolding makes human identities visible, governed and accountable. Such visibility is far less common for identities associated with APIs, service accounts, chatbots and AI agents operating in the same environment.

“The identities you cannot see are the ones that get you breached,” Watts says. “Unlike employees, AI agents arrive silently.” An AI agent wouldn’t have to be malicious to create a problem. It could be deployed without IT oversight, entering your environment via an API key from an SaaS app.

2. What can AI agents connect to?

To perform specific tasks, AI agents might interact with agency databases, APIs, model context protocol (MCP) servers, SaaS platforms and enterprise application environments. “The more access and data an agent has, the more useful it becomes,” Watts says. “But it also introduces more risks.”

Determining what AI agents should be able to access from the start is essential to secure operations.

3. What can AI agents do?

Agents act by reasoning their way around friction to achieve desired outcomes. “They’re thinking through problems and coming up with ways to get around them,” Watts says. As the underlying models grow more capable, that same reasoning can lead agents into actions they were never explicitly permitted to take. Attackers could target agents with weak identity controls and exploit systems. Agencies must have a way to stop an agent instantly when it moves outside its boundaries.

Advanced identity and access management (IAM) is the only solution that answers all three of these questions at the same time.

Governing a New Identity Type

Historically, humans have received “first-class” identity status because they make decisions that were impossible for machines. Now that AI can make similar decisions, agencies face a new paradigm in identity.

“AI agents read and write in your systems, reason through contextual nuances and take action on instructions. That’s not a service account,” Watts says. “That is a first-class identity.”

Key AI governance tools include:



Kill switches. Advanced identity controls revoke agent sessions and access tokens across connected systems when a threat is detected.



Human-in-the-loop approvals. Policy enforcement requires verified human approval before agents perform sensitive or high-risk actions.



Runtime enforcement. AI agents should not rely on broad standing privileges. Instead, policies should evaluate what the agent, user and task are allowed to do at the time of request.



Agent lifecycle management. The best software continuously reviews agent access, revoking it when agents get decommissioned.



Unified identity control. A complex IT environment needs a centralized platform that creates visibility around agents, governs their behavior and provides



transparent audit trails. These platforms include tools to discover shadow agents, register them centrally, enforce least-privilege access and provide instant cross-system revocation.

“If your human identities live in one system, your service accounts are in another and your AI agents are in a third, then you don’t have governance,” Watts says. “You have three identity silos sharing a building. A unified identity platform provides a single control plane across all your identities in one consolidated view.”

Conclusion

The speed of cyberattacks shows the contrast between human and agentic identity threats. If an adversary compromises a human identity, risks accumulate at human speeds over days or weeks. But within seconds, an AI agent identity threat could infiltrate connected tax databases, health records and financial systems.

“A compromised AI agent would cause damage at machine speed, around the clock, across every system it can access,” Watts says.

That doesn’t have to happen. Comprehensive identity controls that treat AI agent identities like human identities give state and local agencies the best opportunity to safely reap the benefits of agentic AI.

1. <https://www.govtech.com/artificial-intelligence/nascio-report-offers-new-data-guidance-around-agentic-ai>

This piece was written and produced by the Government Technology Content Studio, with information and input from Okta.



Produced by Government Technology

Government Technology is about solving problems in state and local government through the smart use of technology. Government Technology is a division of e.Republic, the nation's only media and research company focused exclusively on state and local government and education.

www.govtech.com



Sponsored by Okta

Okta, Inc. is The World's Identity Company™. We secure AI, machine, and human identity so everyone is free to safely use any technology. Our customer and workforce solutions empower businesses and developers to protect their AI agents, users, employees, and partners while driving security, efficiencies, and innovation. Learn why the world's leading brands trust Okta for authentication, authorization, and more at okta.com.