

“Noses In, Fingers Out” Meets AI

By Fayeron Morrison

The universal rule of governance oversight has long been that boards monitored operations and controls, knew who was in charge of what, but would then leave managers to do their jobs. What happens, though, when autonomous AI agents increasingly take action on their own for the company? Or, when employees tap “shadow” AI tools independently to perform tasks? New federal policy now means that you—and your board—can face felony charges for AI misuse.

There is a phrase that has governed boardrooms for decades. You have heard it in director training programs, read it in governance textbooks (and probably nodded along to it in your first board orientation). “Noses in, fingers out.”

The idea is elegant in its simplicity. Boards stay informed, keeping their noses in the business, but resist the temptation to meddle in management decisions, keeping their fingers firmly out of day-to-day operations. It was designed to protect the boundary between governance and management. It gives CEOs operational latitude and prevents boards from micromanaging executives.

“Noses in, fingers out” worked for boards when the most consequential actors inside a company were human—visible, accountable, and replaceable. AI agents break that assumption.

For most of the twentieth century, it worked well. It does not work anymore, though, and artificial intelligence is the reason why.

“Noses in, fingers out” worked for boards when the most consequential actors inside a company were human—visible, accountable, and replaceable. AI agents break that assumption. They do not just generate text for review. They act, accessing data, trig-

gering workflows, making decisions, and operating at a speed no quarterly oversight can match. When something goes wrong, the governance question is immediate: Who authorized this, what was the agent allowed to do, and where is the evidence?

On June 2, 2026, that question moved into the federal enforcement spotlight. Executive Order 14409 directed attention to existing criminal laws when AI is used to unlawfully access systems or data. If your board is discussing AI but cannot produce an agent inventory, scoped access rights, audit trails, vendor controls, and a shadow AI discovery process, it is not governing AI.

In June 2026, Eric Brandwine, a vice president and distinguished engineer at Amazon Security, argued that human-in-the-loop oversight can fail when repeated exposure causes people to stop responding meaningfully. It is a striking observation from a senior executive at one of the world’s largest technology companies, and it reflects the operational reality underlying the governance argument below.

“Noses in, fingers out” was designed for a world where consequential decisions were made by human beings, executives who could be hired, evaluated, and if necessary, fired. A board could reasonably monitor a CEO’s judgment by reviewing quarterly financials, strategy sessions, and management presentations. The information flow was curated and comprehensible.

Artificial intelligence breaks every one of those assumptions. Today, the most consequential decisions inside many organizations are not made by executives at all. They are made by algorithms—credit models, hiring filters, pricing engines, fraud-detection systems, and increasingly, autonomous AI agents authorized to take actions without human review.

These systems operate continuously, at scale, and at speeds no human overseer can match. They make thousands of decisions per hour that individually

Fayeron Morrison CPA, CFE, is founder and president of *Elemental AI*, based in California. [www.linkedin.com/in/fayeronmorrison]

look trivial, but collectively define the company's risk profile, regulatory exposure, and ethical posture.

A board that reviews quarterly reports cannot govern this reality. A board that relies on management to characterize AI risk is experiencing, in the words of one governance analyst, a simulation of reality, not reality itself. The curated boardroom presentation and the actual behavior of a deployed AI system are two very different things, and the gap between them is where governance failures live.

For years, the argument for updating board governance doctrine was philosophical, but now it is legal.

On June 2, 2026, President Trump signed Executive Order 14409, titled "Promoting Advanced Artificial Intelligence Innovation and Security." Coverage primarily focused on the headline—a voluntary framework for AI developers to give the federal government early access to "frontier" models. Buried in Section 4, though, is something more immediately consequential for every board overseeing a company that deploys AI agents.

For the first time, a presidential directive named AI agents as distinct legal actors in a data-access liability framework. This pointed the full force of existing federal criminal statutes directly at AI agent-enabled actions. There is no need to wait for Congress or for a notice-and-comment rulemaking process before the government can enforce existing statutes. The liability clock is already running.

Until now, a company whose AI agent exceeded its authorized access faced civil exposure. Now, the same conduct sits in the DOJ's priority criminal enforcement queue.

Section 4 directs the Attorney General to prioritize enforcement of three existing federal criminal statutes—identity fraud, unauthorized computer access under the Computer Fraud and Abuse Act, and wire fraud. These apply to anyone who uses AI agents to unlawfully access data for criminal purposes.

Routing through existing statutes rather than creating new law means prosecutions can begin immedi-

ately. The EO resolves prior ambiguity by treating an AI agent operating outside its authorized scope as the functional equivalent of an unauthorized human actor—with all the liability that implies.

Until now, a company whose AI agent exceeded its authorized data access faced civil exposure: a lawsuit or regulatory inquiry. After June 2, the same conduct sits in the DOJ's priority enforcement queue. The prospect of criminal enforcement changes the risk calculation considerably.

The SEC has made its priorities equally clear. Its Fiscal Year 2026 Examination Priorities explicitly direct examiners to assess whether firms have adequate policies to monitor and supervise AI technologies and to review whether AI representations are accurate. Passive, noses-in oversight produces no paper trail and fails that test.

For boards with EU exposure, important requirements governing certain high-risk AI systems are now scheduled to apply beginning in December 2027, with other requirements following different implementation dates. Boards that wait until 2027 to act will find themselves in the same position U.S. boards are today—behind, exposed, and scrambling.

Before unpacking the governance implications, it helps to be precise. An AI "agent" is *not* an AI chatbot that answers questions. It is an AI system that takes autonomous actions on behalf of a user or organization. Such actions can include browsing the web, executing code, sending emails, querying databases, making API calls, and initiating transactions. The agent does not just generate a response. It does things.

In March 2026, cybersecurity researcher Jeremiah Fowler discovered three completely unprotected databases belonging to Sears Home Services with no passwords, no encryption. Inside were 3.7 million customer chat transcripts, 1.4 million audio recordings, and nearly four terabytes of data collected by Sears' AI voice agents. The exposed material included names, home addresses, telephone numbers, appliance details, and repair schedules. Some recordings ran up to four hours, capturing background conversations and sounds customers never knew were being recorded.

No hacker or sophisticated attack was involved. The agents had been collecting and storing customer data continuously, at scale, with zero governance controls around what they retained or how it was protected. The systems appear to have collected and stored data as designed, but the surrounding security and governance controls were plainly inadequate.

That is the deployer liability story in its purest form. No attacker, no rogue employee, just an organization that built and deployed AI agents without asking what they were collecting, where it went, or who was watching.

The Sears exposure is a civil and regulatory liability story. Section 4 of the June Executive Order is what happens when that same governance failure moves into criminal territory. In November 2025, Anthropic reported that a threat actor it assessed with high confidence to be a Chinese state-sponsored group had manipulated Claude Code into attempting intrusions against roughly 30 global targets. Security researchers noted that 80 to 90 percent of the attack tactics were carried out by the AI agents themselves, with minimal human involvement.

When an AI agent acts, the action is already taken. A human may not see it until later (if at all). The decision/action cycle that governance frameworks were built around has collapsed.

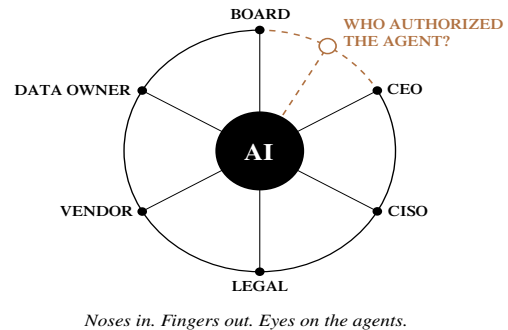
This was not a case of a human hacker using AI as a tool. The agents were the primary actors. This is exactly the scenario Section 4 was written to address, and why AI agent-enabled criminal access is now a federal enforcement priority.

When an AI system generates text, a human reads it and decides what to do. When an AI agent acts, the action is already taken. The human may not see it until later (if at all). The decision/action cycle that governance frameworks were built around has collapsed into a single automated step, and the accountability question that used to have a clear human answer is now much murkier.

That murkiness is precisely what Section 4 is de-

“Noses In, Fingers Out” Is Dead **Who Answers For AI?**

“Noses In, Fingers Out” Is Dead. AI Killed It.



Noses in. Fingers out. Eyes on the agents.

signed to resolve. The answer regulators are arriving at, on both sides of the Atlantic, is the same—the organization that deployed the agent is accountable. Proving accountability now requires documentation that most companies do not currently have.

AI agent liability spans three layers: developer, deployer, and user. Most boards sit in the deployer layer, and that is where the most underappreciated exposure lives. Legally, deployers face liability for agent behaviors that harm others even when the underlying model performed exactly as documented. Your configuration choices (permissions, scope definitions, oversight mechanisms) are your liability. The four items below represent the minimum documentation that prudent deployers should be prepared to produce.

The liability standard the new Executive Order creates is not vague. To demonstrate that an AI agent operated within approved boundaries—and to strengthen the company’s position if its practices are questioned—a company should be able to produce four things:

- *Authenticated agent identity.* Every AI agent must be linked to a named human authorizer. Anonymous or unattributed agents spun up by employees without governance oversight are unauthorized by definition.

- *Operation-level access control.* It is not sufficient to say an agent had access to the database.

Authorization must be scoped at the transaction level: the data the agent was permitted to touch, in what context, and for what purpose. Sears' agents had standing authority to collect and retain far more than any customer consented to or any governance process approved. Broad, general permissions are not a defense, but rather a problem regulators are looking for.

□ *A contemporaneous, tamper-evident, immutable audit trail.* This is where most companies currently fail. While logs exist, logs and audit trails are *not* the same thing. An audit trail is a coherent, time-stamped, tamper-resistant chain of custody for every action an agent took. Such a record can withstand a federal investigation.

A February 2026 study by researchers from MIT, Harvard, Stanford, and Carnegie Mellon documented how AI agents can bypass role-based access controls through conversational prompting alone. Some 61 percent of organizations currently rely on fragmented logs that cannot produce a coherent chain of custody for a single agent interaction.

□ *Encryption that meets federal security standards on all data the agent accessed.* AI agent governance is now being treated as a cybersecurity compliance issue, not just a technology ethics question. If your agents touch sensitive data, the encryption protecting that data needs to meet the same standards required of federal systems. Your CISO should be able to confirm this in writing.

Then there is the “shadow AI” problem. Verizon reported that 45 percent of employees in its dataset regularly used AI tools on corporate devices and that approximately two-thirds accessed those tools through noncorporate accounts. According to a 2026 Verizon report, this shadow AI has become the third most common non-malicious data loss trigger, a fourfold increase year over year.

Such activity may involve no malicious intent and may occur without triggering an immediate breach notification. Just proprietary source code and internal documents quietly moving into systems your board never reviewed and your legal team has never approved.

Shadow AI nevertheless creates a serious gover-

nance problem. An employee may connect an unapproved third-party agent to company systems without understanding what information the agent can retrieve, retain, or transmit. Even when the employee's purpose is benign, the resulting access can violate company policies, contractual obligations, privacy requirements, or data-security controls. If the information is later used unlawfully, the company may also face difficult questions about authorization, monitoring, and accountability.

The governance implication is direct: Passive oversight cannot govern what it cannot see. A quarterly security report may not reveal that an unapproved AI agent has been accessing customer data. Boards therefore need reliable agent-discovery and inventory capabilities. “Noses in, fingers out” was never designed to provide that visibility.

“We have an AI policy” is not a defense. “Here is our documented board-level review of the model’s performance against its risk thresholds last quarter” is.

The Harvard Law School Corporate Governance Blog observed that the most urgent AI governance failure is not ignorance—it is inaction. Most boards are discussing AI, but very few have formal oversight mechanisms with teeth.

This is not a criticism of individual directors, but a structural problem. The “noses in, fingers out” model was never designed to produce the verifiable, documented governance evidence that regulators now expect. Board discussions generate minutes. Minutes are not audit trails of AI risk oversight. Regulators now will want evidence: model cards, explainability documentation, incident logs, proof that the board reviewed and challenged outputs.

“We have an AI policy” is not a defense. “Here is our documented board-level review of the model’s performance against its risk thresholds last quarter” is.

None of this is an argument for boards to start running operations. The legitimate concern behind “fingers out”—that board micromanagement under-

mines executive authority—remains valid. Instead, the answer is to redefine what informed oversight requires in an AI-enabled structure.

In practice, boards now need documented answers to five AI questions before the next examination cycle:

□ *AI agent inventory*. Every deployed agent, named, with its authorizing human owner, its data access scope, and its operational authority documented.

□ *Shadow AI discovery*. What is the technical and governance mechanism for identifying agents deployed outside sanctioned channels?

□ *Vendor contract provisions*. Do your contracts include kill-switch rights, compliance warranties, and audit provisions? If not, when will they be renegotiated?

□ *Audit trail integrity*. What does your audit trail look like for the last ninety days of agent activity? Not logs, but a coherent, producible chain of custody.

□ *SEC readiness*. If the examiner arrived tomorrow, what would you show them? The answer to that question is your actual governance posture.

This governance gap also raises fiduciary concerns. Directors are expected to make informed decisions and, under Delaware oversight doctrine, to make good-faith efforts to establish and monitor systems for reporting material risks. As AI agents assume

more consequential responsibilities, boards should consider whether existing reporting systems provide meaningful visibility into those risks.

When an AI agent takes a consequential action—moving money, denying a claim, or accessing sensitive data—and something goes wrong, one of the first questions will be who authorized it. The answer should trace to a human being with documented authority who made a deliberate governance decision to deploy the agent with that scope and permit it to take that category of action.

If you are not sure where your company stands, you should find out. Start with the self-assessment at our company website [<https://elementalai.ai>].

The old doctrine “noses in, fingers out” gave boards a simple rule for a simpler world. It assumed the most consequential actors inside your organization were human, visible, accountable, and replaceable. That world is gone. The agents are acting now, continuously, at scale, in your name.

The new doctrine is simple too: “Noses in, Fingers out, Eyes on the Agents.”

The board that cannot see what its agents are doing cannot govern what they do next. The agent accountability era has a governance infrastructure requirement. ■