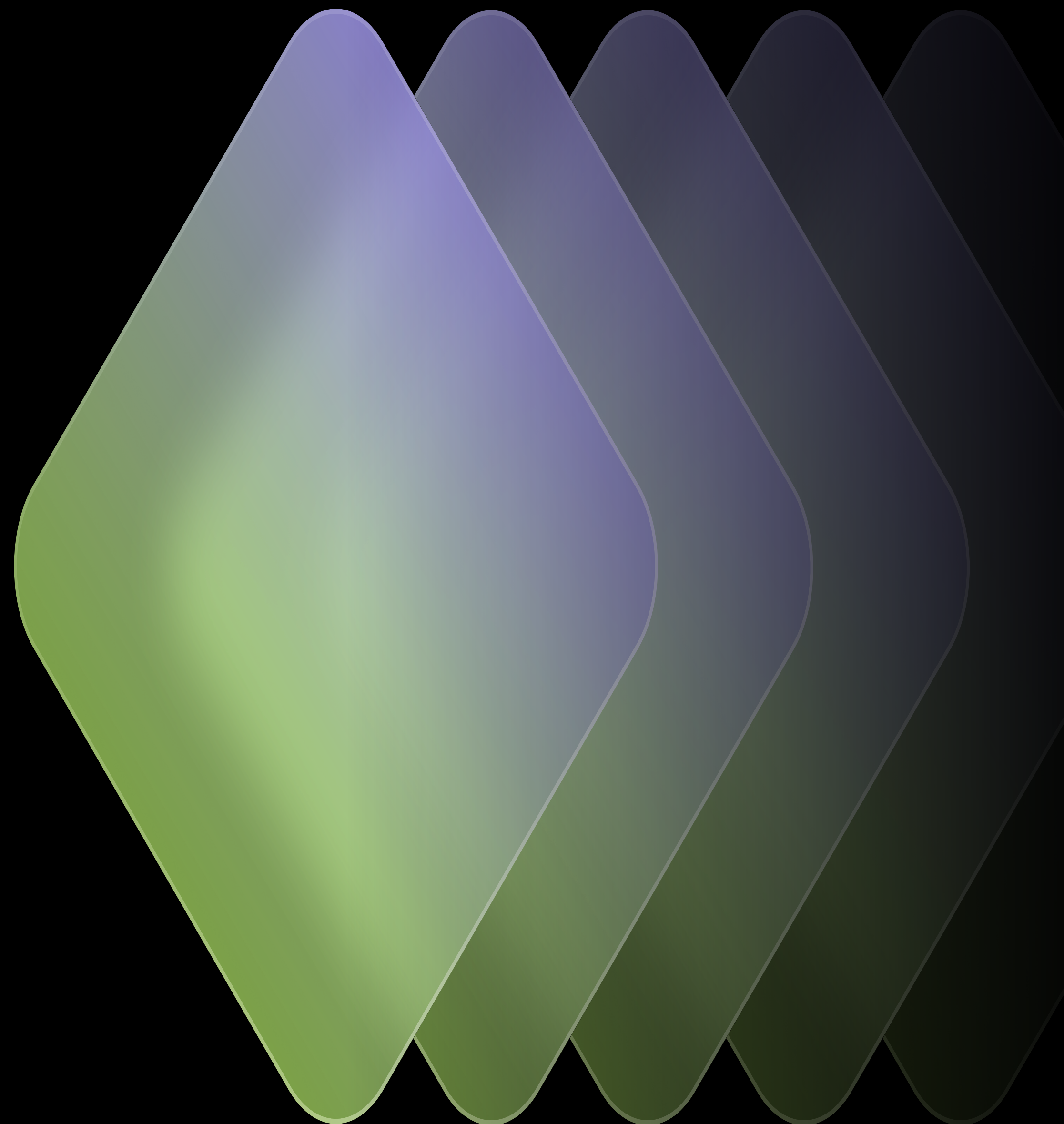


Guide to Deploying AI You Can Actually Trust

2026



We asked IT Directors, Heads of IT, and CIOs at mid-market and enterprise tech companies about their experience and concerns implementing AI. The same worries kept surfacing: misleading output, messy data, security risks, and where human verification fits in.

In this guide, we uncover what stands behind each of these concerns and ways to adopt AI safely. Jump to the section that matters most to you, or read straight through.

Contents

- 01 Weak Access Control and Information Leakage 4
- 02 Potential Security Risks 7
- 03 Misleading Output and Hallucinations 10
- 04 Low Data Readiness 13
- 05 Automation vs. Human-in-the-Loop Balance 16
- 06 Verification Tax Trap 19

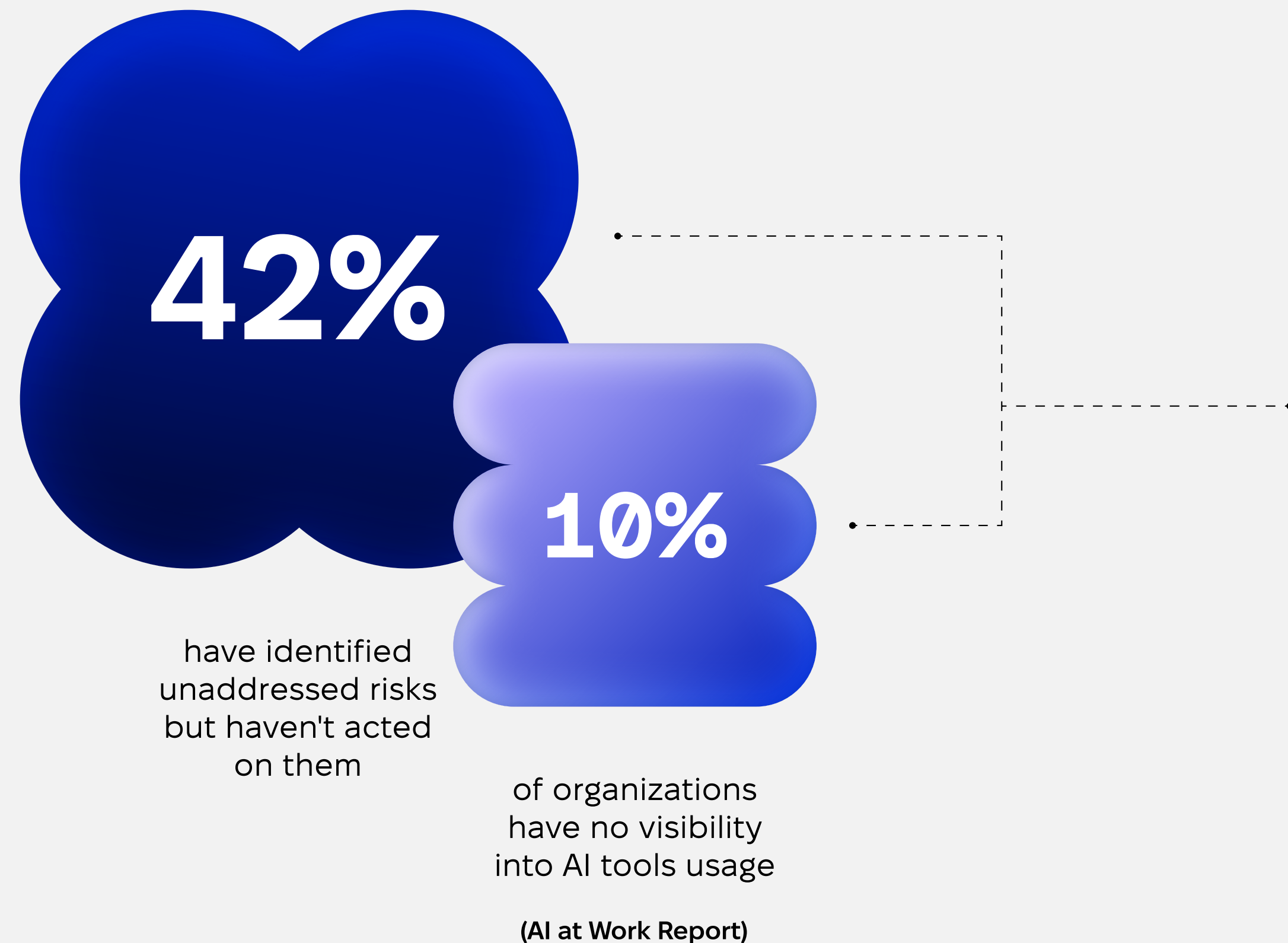
01

Weak Access Control and Information Leakage

"Right now we are still stuck in policy and governance phase. We have teams that want to use AI, but due to data sensitivity and lack of IT support, we struggle to allow the use of AI. I'm still waiting to get my AI policy pushed through Legal."

— Kent, IT Manager / Information Security Officer, Enterprise

Weak Access Control and Information Leakage



What's behind it

AI tools get adopted faster than governance policies. By the time IT catches up, the tool is embedded in daily workflows, already connecting to internal documents that may contain sensitive data. According to IBM's Cost of a Data Breach Report, one in five organizations has faced a breach due to shadow AI, yet only 37% have policies to manage or detect it.

The core gap is permission awareness. Most AI tools retrieve documents based on semantic similarity, not on whether the user is actually authorized to see them. Exposure can happen through legitimate usage, like an employee asking a regular question and accidentally getting access to confidential information.

Risks like that delay AI implementation. Companies often get stuck between the need to optimize and the reality, in which managers evaluate AI tools for months without moving further.

1. Implement RBAC retrieval

For confidential data, reveal only the information that the user is already authorized to see. You can do this through user-level access tokens. Rather than indexing an entire knowledge base upfront, the system authenticates each request as that specific user and retrieves only what they are permitted to access. Since sensitive content is never stored in the index, the system cannot leak data it never held.

2. Control data flows, not tools

Decide which data an AI tool is allowed to touch based on data class and permission scope, not the product name or vendor claims. Build that rule once, apply it to every tool. There's no point in tracking every new feature each vendor ships; you'll always be behind.

3. Move from static permissions to continuous evaluation

Set behavioral baselines and monitor patterns like access frequency, data types, and request volume. Trigger verification or automatic session termination when something unusual happens. While static permissions tell you who was allowed access at setup, behavioral monitoring tells you what's actually happening.

4. Treat the browser as a separate control surface

Enforce browser policies on managed devices, including extension allowlists, approved AI tools, and data classification rules, to manage what employees can upload. Browser extensions are a common and underestimated leakage point. Employees unknowingly upload sensitive information, either to speed up work or test a new tool.

5. Have an AI governance board

Make sure security, legal, IT, and relevant business units have shared visibility and accountability over AI implementation. A board creates a consistent decision-making process for evaluating new tools, developing policies, and handling exceptions, which particularly matters if you handle sensitive customer data and need to achieve regulatory compliance. KPMG's AI governance principles are a useful reference for boards getting started.