

AI Governance Letter · Issue #2

Permission Is Not Practice

Shadow AI and the limits of prohibition — why restriction leaves boards with less information rather than less risk

August 2026 · Reto Gruenenfelder, IT AI Advisory

Impressum

IT AI Advisory Sdn. Bhd. (202601023072)
B-3A-8, Icon City, Jalan SS8/39, 47300 Petaling Jaya, Selangor, Malaysia
Responsible for content: Reto Gruenenfelder
reto@itaiadvisory.com · itaiadvisory.com

Published for informational purposes only. Not legal, regulatory, or compliance advice. Full disclaimer: itaiadvisory.com/disclaimer

Dear Board Director,

When I visit an organisation, I ask people a simple question: do you use AI in your work?

The answer is almost always the same, and it is almost never an answer. People tell me: “We are only allowed to use Copilot.”

I have heard this in banks, in listed companies, and in professional firms. It is offered helpfully and without hesitation. And it should give any board pause, because the question was about what people do, and the reply was about what they are permitted to do. Those are not the same thing, and the difference is where governance quietly fails.

That difference has a name. Shadow AI is the use of AI tools, features and services inside an organisation without approval, inventory or oversight. This letter is about why prohibition tends to widen that gap rather than close it — and why a board that responds to AI with a restriction may end up knowing less than one that does not.

The Answer That Is Not an Answer

“We are only allowed to use Copilot” is a statement about policy. It is not a statement about behaviour. It does not say what tools are open in a browser tab, what is being used on a personal phone during a deadline, or what a colleague pasted into a free chatbot last Thursday.

People answer this way because they have learned to hear the question as an audit. In an organisation where using anything else is a breach of policy, the only safe reply is the permitted one — whether or not it is complete. Nobody is being dishonest. They are answering the question they believe is being asked.

The same substitution happens one level up. When a board asks whether AI use is under control, management can answer, entirely truthfully, that staff are permitted to use only one approved tool. The board hears an assurance about practice. It has received a description of policy. The gap between the two is never discussed, because nobody notices it opened.

Permission is not practice. A board that accepts the first as evidence of the second has not been misled. It has simply stopped asking.

What the Permitted Tool Actually Measures

There is now data on this. Recon Analytics, drawing on more than 150,000 respondents between July 2025 and January 2026, examined what happens when employers provide AI tools with and without alternatives.

Where Copilot is the only platform an employer provides, 68% of workers adopt it as their primary tool. Where both Copilot and ChatGPT are available, Copilot’s share falls to 18%. Where all three major platforms are available, 8% choose Copilot.

The same tool, the same people, three very different numbers. What changes is not the software but the presence of a choice. An adoption figure of 68%, presented to a board as

evidence that the approved tool is working, may be measuring nothing more than the absence of alternatives.

Two further findings from the same research are worth noting. Among workers who tried Copilot and stopped using it, 44.2% cited distrust of its answers — the highest figure of the three major platforms. And more than half of those who have Copilot available at work also have ChatGPT available, which suggests that “only Copilot” is frequently a statement of policy rather than a description of the environment.

A restriction that channels people toward the tool they trust least, in an environment where alternatives are a browser tab away, is not a stable control. It is a queue forming behind a door that does not lock.

One caveat on this evidence: the study covers paid AI subscribers in the United States, a self-selected population, and does not measure how common such restrictions are across enterprises generally or in this region. It shows what restriction produces, not how widespread restriction is. US adoption patterns have, however, tended to precede those elsewhere, and boards in this region may reasonably treat these figures as an indication of where their own organisations are heading.

A Ban Is a Control That Produces No Evidence

Prohibition is an attractive response to an unfamiliar risk. It is quick, it costs nothing, it can be minuted, and it signals seriousness. For a board under pressure to demonstrate that AI is being taken seriously, a restrictive policy is the cheapest available proof.

It is also, in governance terms, close to worthless — because it generates nothing a board can review. There is no inventory of what is in use, no log of incidents, no record of which business processes now depend on a model, and no data on where information is travelling. A prohibition cannot be reported on. It can only be asserted.

This is worth stating plainly: the control removed the evidence. Not the risk. An organisation that has banned AI tools does not have less AI in use than one that has not. It has less knowledge of the AI in use. The board has purchased the appearance of control at the price of the information required to exercise it.

Three Ways Prohibition Increases the Exposure

Beyond the loss of visibility, a ban tends to make the underlying risk worse in three specific ways. Each is worth testing against your own organisation.

1. It moves usage to weaker terms, not to zero

A sanctioned enterprise deployment normally carries a data processing agreement, contractual assurance that inputs are not used for training, single sign-on, audit logging and retention controls. Unsanctioned use carries none of these. It happens in free consumer accounts, often on personal devices, under standard consumer terms. The prohibition rarely eliminates the activity. It strips the protections from it.

2. It is broken first by your own procurement

While staff are told that AI tools are not permitted, AI features arrive inside software the organisation already licenses — an assistant enabled by default in a productivity suite,

summarisation added to a service desk in a routine update, scoring introduced into a recruitment platform. Nobody approved these as AI systems, because nobody purchased them as AI systems. The policy is breached not by employees but by the organisation's own contracts.

This produces an uncomfortable position. The single permitted tool is frequently the one that received the least scrutiny, precisely because it arrived through an existing vendor relationship rather than through a procurement decision. The organisation has restricted everything it examined and permitted the thing it did not.

3. It silences the people best placed to raise the alarm

If a colleague notices that an AI system has produced a confidently wrong answer, a biased output, or a plausible fabrication, they can normally only report it by disclosing that they were using the tool. Where that is a disciplinary matter, the report does not arrive. The failures a board most needs to hear about are precisely those that remain unspoken.

What the Breach Data Suggests

IBM's 2026 Cost of a Data Breach Report, published on 29 July 2026 and based on breaches at 602 organisations between March 2025 and February 2026, offers a useful corrective to the assumption that AI risk is principally about models.

More than one in five organisations reported a breach targeting AI models or applications. The most common causes were not the models themselves. They were weaknesses in the surrounding systems: compromised APIs, applications or plug-ins (27%), and cloud misconfigurations affecting AI workloads (27%).

That distribution matters for this discussion. The exposure sits in the connective tissue around AI — the integrations, extensions and configurations that accumulate without a procurement decision and are rarely inventoried. A policy restricting which chatbot staff may open does not address any of it.

Follow-on research by the Ponemon Institute, conducted in May 2026 with 456 of those organisations, contains a finding that deserves a moment's reflection. 85% said they planned to increase security spending after becoming aware of advanced frontier AI capabilities, compared with only 64% who said the same after actually experiencing a breach. Anticipated risk moved these organisations more than experienced risk did. Where a prohibition has removed the capacity to detect incidents, there may be very little experienced risk to learn from.

Questions Worth Raising at Board Level

The following are framed so that they cannot be satisfied by a statement of policy. Each asks for evidence rather than permission.

1. What AI is actually in use, as distinct from what is permitted?

A list of approved tools is not an answer. The question is what management can demonstrate about actual usage, and by what method it was established.

2. Which AI features arrived through vendors rather than through a procurement decision?

Assistants and AI functions enabled by default inside existing licences are the most common unassessed AI in most organisations. They will not appear on a list of purchases.

3. If a member of staff found an AI output to be wrong, how would we hear about it?

If the only route requires admitting to a policy breach, the organisation has no reporting channel, irrespective of what the incident process says on paper.

4. What evidence does our AI policy generate?

Controls that produce no data cannot be monitored. It may be worth asking what the board would expect to see if the policy were being widely ignored — and whether it would look any different.

5. Was our permitted tool assessed, or merely permitted?

Was there a data flow review, a decision on training and retention, and a named owner? Or was it approved because it was already present in the environment?

6. Where is our data going, under whose terms?

Sanctioned deployments and consumer accounts carry materially different contractual protections. The board should know which one the organisation is relying on.

A Final Observation

This pattern is not new. Organisations met consumer file-sharing and personal devices with prohibition a decade ago, and learned that restriction without a credible alternative produces invisible adoption rather than no adoption. Those that moved quickest to provide a sanctioned, genuinely usable option regained the visibility they had lost. Shadow AI is the same problem with a shorter fuse and a wider blast radius.

None of this argues for permissiveness. It argues for controls that produce evidence. A restriction that makes a risk invisible has not reduced it — it has removed the board's ability to see whether it is growing.

The question is not whether your organisation has a policy on AI. It is whether that policy tells you anything you could take to a board meeting.

About the Author

Reto Gruenenfelder advises executive and non-executive boards on AI governance through IT AI Advisory, drawing on a 20-year career at IBM Financial Services and his former role as Technology Advisor to the Board of Bursa Malaysia.

Primary Source Documents

The following links directly to the original material referenced above.

→ [**IBM: One in Four Malicious Breaches are AI-Enabled — 29 July 2026**](#)

→ [**IBM: Cost of a Data Breach Report 2026 — full report**](#)

→ [**Recon Analytics: AI Choice 2026 — Why Licenses Don't Equal Adoption — February 2026**](#)

Published for informational purposes only. Not legal, regulatory, or compliance advice.

Full disclaimer: itaiadvisory.com/disclaimer

© 2026 IT AI Advisory Sdn. Bhd. (202601023072). This letter may be shared freely in its original, unmodified form provided that itaiadvisory.com is identified as the source.