
AI Risk and the Availability Bias Problem

Managing AI risk seriously while keeping it in proportion to the risks that cause most real-world data exposure

By Jeff Franzen - jbf.com - June 22, 2026

Executive Summary

Daniel Kahneman's research teaches that humans routinely give more weight to risks that are vivid, novel, or easy to recall, while familiar risks can receive less attention even when they occur more often. AI is a real data-security concern, but credential theft, phishing, sharing mistakes, lost devices, exposed systems, and uncontrolled exports remain frequent causes of organizational data exposure. AI belongs on the risk register, but good risk management requires prioritizing threats according to actual likelihood and impact rather than visibility and public attention.

We should manage AI risk seriously, and we should prioritize risk proportionally. The challenge is not that organizations will ignore AI; the challenge is that AI can absorb so much institutional attention that more common sources of data loss receive less sustained focus.

01

Core Claim

AI creates real data-security risks. Those risks deserve governance, training, monitoring, and policy.

But AI should not automatically be treated as the largest organizational data risk simply because it is new, visible, and heavily discussed. Kahneman's work on the availability heuristic reminds us that humans often give added weight to risks that are dramatic or easy to imagine, while underweighting risks that are familiar, routine, and statistically more common.

Good risk management should be proportional. It should prioritize threats based on likelihood and impact, not visibility alone.

The Availability Bias Problem

Kahneman showed that people often judge the probability of an event by how easily examples come to mind.

That is why people may focus heavily on terrorism, shark attacks, or plane crashes while giving less attention to more common risks such as falls, vehicle accidents, chronic health issues, or ordinary workplace hazards.

The same pattern can occur with AI. Because AI is new, powerful, and constantly discussed, it can feel like the dominant threat to organizational data. But visibility is not the same thing as probability.

Where Data Exposure Usually Happens

Organizations lose or expose sensitive data through many familiar pathways. AI belongs on this list. It deserves attention. It should be managed as part of the broader data-exposure program.

Risk Category	Examples
Exposed systems	Open ports, weak patching, public services, unreviewed configurations
Credentials	Stolen passwords, weak MFA, account takeover, reused credentials
Phishing	Social engineering, email compromise, malicious links and attachments
Cloud sharing	SharePoint, ArcGIS, Power BI, and external-link mistakes
Devices	Lost laptops, phones, removable drives, and unmanaged endpoints
Documents	Sensitive files stored in the wrong locations or retained too long
Exports	Spreadsheets, CSV files, offline extracts, and uncontrolled copies
AI tools	Unapproved tools receiving sensitive information or organizational data

The Better Risk Question

The wrong question is: "What risk is most visible right now?"

The better question is: "What risks actually cause the most incidents, the greatest losses, or the highest probability of harm?"

That is the difference between perceived risk and actual risk. A mature organization allocates resources based on evidence, not headlines.

Red Cross Example

If advising the American Red Cross, I would recommend continuing to build AI guardrails, employee education, approved-use policies, and monitoring for unapproved AI tools.

I would also recommend keeping AI in context with other recurring exposure paths. A single compromised password, phishing attack, accidentally shared SharePoint folder, exposed ArcGIS service, or uncontrolled spreadsheet export can expose more sensitive information than many AI scenarios.

The point is not to minimize AI risk. The point is to place AI risk in its proper context.

Practical Framework

A better organizational framework treats AI as one category inside a broader data-exposure program. AI governance should be integrated into the existing risk program, not isolated as a separate category that draws attention away from other controls.

The goal is not to argue that AI is harmless. The goal is to manage it with the same discipline used for other sources of organizational exposure: inventory the risk, classify the data, set rules, monitor behavior, train people, and measure outcomes.

Closing Argument

AI is on the list. It deserves governance. It deserves training. It deserves policy.

But if our objective is to reduce organizational risk, focusing exclusively on AI while giving less attention to the more common causes of data exposure would miss the central lesson of Kahneman's work.

The responsible position is neither dismissal nor overcorrection. It is proportionality.