



FIELD NOTE

Adopting AI without the governance bureaucracy

A practical take on letting your team use AI fast, with the few guardrails that actually matter.

A field note from AEGYS
aegys.io · secure@aegys.io

"What is our AI governance strategy?" is a question that has stalled more momentum than any actual AI risk. You can be genuinely responsible about AI without standing up a committee, a 40-page policy, or a six-month framework rollout. Responsible and bureaucratic are not the same thing, and at mid-market scale, the bureaucracy is the bigger threat to getting any value at all.

Borrow the structure, skip the ceremony

The good frameworks already separate the signal from the ceremony. NIST's AI Risk Management Framework organizes the whole problem into four plain verbs: **Govern** (decide who is accountable), **Map** (understand each use and its risks), **Measure** (check that it works and watch for drift), and **Manage** (treat the biggest risks, keep a human in the loop). That is not a binder. For a mid-market team it is four habits, and you can adopt them this week without anyone forming a working group.

The guardrail fits on one page

Most "AI policy" effort goes into a document nobody reads. Replace it with one page that sorts AI use into three lanes: green for low-sensitivity work (go), yellow for internal or customer data (use an approved tool, ask first), and red for regulated data and secrets (never into public tools). Name the approved tools and the setting to use, give people a human to ask, and you are done. You can tighten it later; you will have something real in week one instead of a perfect document in month six.

Know the short list of things that actually go wrong

You do not need to fear everything; you need to handle a handful of failure modes. The OWASP Top 10 for LLM Applications is the practitioner's list, and the ones a mid-market team meets first are few: **prompt injection** (LLM01), where content the model reads hijacks what it does; **sensitive information disclosure** (LLM02); and **excessive agency** (LLM06), where an AI tool is handed more reach than the task needs. The defenses are unglamorous and familiar: treat model input as untrusted, keep regulated data out, and give tools least privilege.

Ship a win, then widen

Frameworks describe; momentum convinces. Pick one boring, high-volume task where a mistake is cheap and a human still reviews the output, and let AI take the first draft. Measure the time saved. That single visible win does more to build a responsible AI culture than any policy, because it shows the safe path and the fast path are the same path. If you want the formal badge later, a standard like ISO/IEC 42001 exists for AI management systems, but earn the habits first; the certificate should describe a practice you already run.

The thesis

Governance is a means, not the product. Keep the few guardrails that matter, drop the theater, and spend the saved energy shipping safe, useful AI.

Our Practical AI Playbook is the do-it-yourself version of this. If you would rather have a partner run it with you, start with a Security and AI Posture Assessment. Reach us at secure@aegys.io.

References

NIST Artificial Intelligence Risk Management Framework (AI RMF 1.0), NIST AI 100-1.
<https://doi.org/10.6028/NIST.AI.100-1>

OWASP Top 10 for LLM Applications (2025): LLM01 Prompt Injection; LLM02 Sensitive Information Disclosure; LLM06 Excessive Agency. <https://genai.owasp.org/llm-top-10/>

ISO/IEC 42001:2023, Artificial Intelligence Management System. <https://www.iso.org/standard/81230.html>