



Enterprise AI Risk Checklist

Seven risks · warning signs ·
readiness score

How to use this checklist

Work through the seven risks with your team. For each risk, tick every warning sign that is true of your organization today — each ticked box is a flag, so fewer is better. Count the boxes you tick per risk, then total them and read your readiness score at the end.

Enterprise AI doesn't usually fail because the technology doesn't work — it fails on risks that were predictable and manageable. De-risking is not caution; it is the fastest reliable path to production.

Risk 1 — Pilot purgatory · Plan

Pilots that impress in a demo and never reach production or measurable value — the single most common failure.

Warning signs to audit

- ☐ Selection driven by the demo ("look what it can do") rather than production criteria
- ☐ No named business owner accountable for the outcome
- ☐ No success metric tied to a measurable baseline
- ☐ The pilot is run by a lab or innovation team with no path into operations
- ☐ "We'll figure out integration later"

The move: Redesign the workflow before building, define production and kill criteria at scoping, pilot the highest-value use case with a named owner and a measurable target — and be willing to kill pilots that don't clear the bar.

Boxes checked for this risk: ____ 0 = strong · 1–2 = watch · 3+ = exposed

Risk 2 — Data readiness · Plan / Build

AI is only as good as the data beneath it, and most enterprise data is fragmented, siloed, and not AI-ready.

Warning signs to audit

- ☐ "We'll clean the data later"
- ☐ No data lineage or ownership
- ☐ The same field means different things in different systems
- ☐ The AI team has no access to the systems where the real data lives
- ☐ Nobody has budgeted for data preparation

The move: Assess data readiness before choosing a model, establish governance (provenance, quality, access) early, budget 20–30% of the effort for data work, and design a governed access layer where data can't leave its boundary.

Boxes checked for this risk: ____ 0 = strong · 1–2 = watch · 3+ = exposed

Risk 3 — The expanding attack surface · Build

AI, and agents especially, opens a new and non-deterministic attack surface — prompt injection, jailbreaks, exfiltration, over-permissioned agents.

Warning signs to audit

- ☐ Agents authenticated with shared API keys rather than distinct identities
- ☐ No red-teaming or adversarial testing before launch
- ☐ Broad, standing permissions instead of least-privilege
- ☐ No visibility into which AI tools employees actually use
- ☐ "The model won't do anything bad — we tested it once"

The move: Red-team before production, give each agent a distinct least-privilege identity, put guardrails and confidence thresholds in front of consequential actions, and govern shadow AI with a sanctioned path rather than a ban.

Boxes checked for this risk: _____ 0 = strong · 1–2 = watch · 3+ = exposed

Risk 4 — Silent model degradation · Run

Models degrade the second they go live — data drifts, the world changes, and providers update models underneath you.

Warning signs to audit

- ☐ Ship-and-forget delivery with no monitoring plan
- ☐ No evaluation set tied to releases
- ☐ No alerting on accuracy, drift, or hallucination rate
- ☐ No named owner for the model after go-live
- ☐ A hosted-model dependency with no plan for provider-side changes

The move: Treat evaluation and monitoring as architecture: eval harnesses in CI/CD, drift detection, alerting, and a named owner for the live system — with model versioning and a rollback path.

Boxes checked for this risk: _____ 0 = strong · 1–2 = watch · 3+ = exposed

Risk 5 — Regulatory exposure · Plan / all

Regulation is a live constraint now. The EU AI Act is extraterritorial, and adding AI to a regulated system can trigger high-risk classification.

Warning signs to audit

- ☐ No one has classified your AI systems by risk tier
- ☐ A plan to fine-tune a model with no awareness that substantial modification can change your role from deployer to provider
- ☐ No logging or documentation designed for auditability
- ☐ No view of which frameworks (EU AI Act, NIST AI RMF, ISO/IEC 42001) apply to your markets

☐ "We'll deal with compliance if it becomes an issue"

The move: Classify every AI system by risk tier at scoping, design compliance in (logging, documentation, human oversight as build artifacts), and get a deployer-versus-provider read on your architecture before you fine-tune. Not legal advice.

Boxes checked for this risk: _____ 0 = strong · 1–2 = watch · 3+ = exposed

Risk 6 — Vendor lock-in and accountability gaps · Plan / Build

AI adoption creates deep dependencies on model providers, platforms, and partners — and outsourcing development does not outsource accountability.

Warning signs to audit

- ☐ Architecture wired to a single model with no abstraction layer
- ☐ No exit path or data-portability plan
- ☐ A vendor who can't show real, in-house AI depth, or whose case studies are thin
- ☐ Unclear IP and model or data ownership terms
- ☐ No read on whether the vendor's approach changes your own regulatory role

The move: Prefer model-agnostic architecture with a fallback path, keep IP and data ownership explicit and in writing, vet a partner's real depth and governance before signing, and retain the ability to move.

Boxes checked for this risk: _____ 0 = strong · 1–2 = watch · 3+ = exposed

Risk 7 — Ungoverned agentic autonomy · Run / all

Agents are reaching production faster than governance can follow — autonomous action at machine speed, across systems and permissions, without clear accountability.

Warning signs to audit

- ☐ Agents with broad, standing permissions and no kill switch
- ☐ No human in the loop on consequential actions
- ☐ Unclear accountability — leadership can't say who is responsible when an agent acts wrongly
- ☐ No audit trail of agent decisions
- ☐ "Pilot" agents quietly wired into production workflows without governance

The move: Define where a human must stay in control, put confidence thresholds, approval gates, and a real kill switch around agent actions, give agents scoped identities and audited permissions, and adopt an AgentOps practice.

Boxes checked for this risk: _____ 0 = strong · 1–2 = watch · 3+ = exposed

Your readiness score

Add up every box you ticked across all seven risks (35 possible). Fewer boxes means a stronger, more production-ready posture. Find your total in the table below.

Total boxes ticked: _____ / 35

Score	What it means
0 – 6	Strong posture. You're de-risking well — keep monitoring the risks you flagged and revisit quarterly.
7 – 15	Manageable gaps. Fix the risks where you ticked 3 or more boxes first; most are upstream, Plan-phase issues.
16 – 25	Elevated exposure. Start with the Plan-phase risks (pilot purgatory, data readiness, regulatory) before scaling further.
26 – 35	High exposure. Treat this as a stop-and-fix before scaling — the failure modes here are predictable and costly.

Reading it by risk

Any single risk where you ticked 3 or more boxes deserves attention regardless of your total — a concentrated weakness in one risk can sink a program on its own. Use the phase tags to sequence the work: fix Plan risks before you build, wire Build risks in during development, and stand up Run risks before anything reaches production.

Turn your score into a plan: bring your AI initiatives to SumatoSoft and our engineers will map them against these seven risks, flag where you're most exposed, and give you a de-risking plan that fits your delivery timeline. Schedule a readiness call at sumatosoft.com/services/ai-readiness-assessment.

ISO 27001 · ISO 9001 · 350+ products · 14+ years · 25+ countries · 98% satisfaction. Practitioner guidance, not legal advice; companion article at sumatosoft.com/blog/enterprise-ai-adoption-risks.



Thank you for your time!

Any questions? Drop us a line!

Headquarters

One Boston Place, Suite 2602
Boston, MA 02108, United States

Other ways to get in touch

info@sumatosoft.com
sumatosoft.com