

An AI policy your audit committee can sign.

*A starting-point policy and risk-appetite statement,
written to survive a committee review.*

HOW TO USE THIS TEMPLATE

An AI policy your audit committee can sign

A starting-point policy and risk-appetite statement, written to survive a committee review.

Somewhere in your company, AI is already touching a number that will end up in front of your board. The question your audit committee will eventually ask is simple: who approved that, and under what policy?

This template gives you the policy. It adapts the structure of COSO's 2026 guidance, *Achieving Effective Internal Control Over Generative AI*, and COSO's Enterprise Risk Management framework into a document a public-company audit committee can review, mark up, and sign. Part A is the policy. Part B is the risk-appetite statement. Part C is the sign-off page, with the questions your committee should ask before anyone signs.

Three rules before you touch it:

- **Every [bracketed] item is yours to fill.** Names, thresholds, systems, tiers. A policy with someone else's placeholders still in it will not survive a committee review, and it shouldn't.
- **This is education, not legal or audit advice.** Have counsel review it against your jurisdiction and disclosure obligations, and walk it past your external auditor before you rely on it for anything touching financial reporting.
- **Shorter beats longer.** Strike any section that doesn't apply to how your company actually uses AI. A committee signs what it understands.

PART A

[Company] Generative AI Use and Governance Policy

Status: DRAFT for audit committee review · Version [X.X] · Owner: [Name, Title] · Effective date: [date]

1. Purpose

This policy governs the use of generative artificial intelligence (GenAI) across [Company], with heightened requirements where AI touches financial reporting, forecasting, disclosures, or other information relied on by the board, lenders, investors, auditors, or regulators. Its objective is to let [Company] capture the efficiency and insight GenAI offers while keeping management able to stand behind every number and statement the technology touches.

This policy adopts the structure of the COSO Internal Control — Integrated Framework as applied to GenAI in COSO's 2026 publication, *Achieving Effective Internal Control Over Generative AI*. Where this policy and that guidance diverge, the divergence is deliberate and documented in Appendix [X].

2. Scope and definitions

This policy applies to all employees, contractors, and third parties acting on [Company]'s behalf, and to all GenAI use — sanctioned or not — that touches [Company] data, systems, or work product. Consistent with COSO's 2026 guidance, GenAI is understood here as machine learning capable of creating new content (text, images, code, analyses), including AI agents that plan and execute multi-step tasks. Traditional rule-based automation and RPA are governed under [existing IT policy reference], not this policy.

Shadow AI. COSO defines shadow AI as unauthorized or ungoverned AI operating outside formal IT oversight. Use of personal or unapproved AI tools for [Company] work is within scope of this policy: a browser tab counts. Discovery of shadow AI triggers the inventory process in Section 5, not automatic discipline, because this policy's first goal is visibility.

3. Governance and accountability

3.1 Board and committee oversight. The audit committee receives reporting at least [quarterly] on GenAI adoption, key risk indicators, incidents, and material changes to high-impact use cases. Consistent with COSO's guidance on oversight, the committee has the standing authority to challenge assumptions, request independent validation, and direct corrective action.

3.2 AI governance committee. [Company] maintains a cross-functional AI governance committee — at minimum legal, compliance, IT/security, risk, and finance — that meets [monthly/quarterly] to approve high-impact use cases, review risk assessments, and maintain this policy. Chair: [Name, Title].

3.3 Named ownership. Every GenAI tool and use case has one named owner accountable for its objectives, risk profile, and controls. Ownership extends to the configuration elements COSO identifies as requiring the same

rigor as any controlled system setting: prompts, system prompts, retrieval data sources, and transformation rules. No owner, no deployment.

3.4 Executive accountability. A single senior executive, [Name, Title], is accountable for AI-related risk across the enterprise and is the escalation point for the committee.

4. Acceptable use

Consistent with the acceptable-use elements in COSO's 2026 guidance, the following apply to all GenAI use at [Company]:

- **Prohibited data.** No personal information, health data, material non-public information, or [other regulated categories] may be entered into any AI tool that has not been approved in writing with data-retention and training settings verified. [List approved tools and the approval owner.]
- **Prohibited uses.** GenAI may not be used for decisions where the law restricts automated decision-making or where fairness exposure is high — including employment decisions — nor to reproduce copyrighted material. [Adapt to jurisdiction with counsel.]
- **Transparency.** AI-assisted work product that leaves the company, or that management relies on, is identified as such per Section 7 and Section 9. Unverified AI output is never represented as verified.
- **Human accountability.** An AI tool is never the reviewer of record. A named person signs for every output that matters, and that person must be able to explain the output without the tool.

5. Use-case inventory and classification

[Company] maintains a living inventory of every GenAI use case, including periodic scans for shadow AI. Each entry records owner, objective, data sources, model and deployment type, and criticality. Each use case is classified using the eight capability types in COSO's 2026 guidance — data ingestion and extraction; transformation and integration; automated transaction processing and reconciliation; workflow orchestration; judgment, forecasting and insight generation; AI-powered monitoring; knowledge retrieval and summarization; and human–AI collaboration — because risk and required controls differ by what the AI actually does.

Each use case is also tiered by where its output lands:

- **Tier 1 — Financial reporting and disclosure.** Output can reach the general ledger, the close, the forecast, or anything the board, a lender, an auditor, or the public will see. Full control requirements of Sections 6 and 7 apply.
- **Tier 2 — Operational decisions.** Output informs internal decisions of consequence but does not enter financial reporting. Sections 6.1 through 6.5 apply.
- **Tier 3 — Productivity and drafting.** Output is fully reworked by a human before any reliance. Acceptable-use rules and training requirements apply.

6. Control requirements

6.1 Outputs are claims, not facts. COSO's first foundational characteristic of GenAI internal control is that GenAI is probabilistic and can be confidently wrong. All controls in this policy treat AI output as an assertion requiring evidence, with human corroboration proportionate to risk.

6.2 Drafting versus deciding. AI drafts; a named human decides. Nothing posts to the general ledger without human approval. Judgment calls — estimates, accruals, reserves, revenue recognition, disclosure language — remain human decisions in all tiers.

6.3 Segregation of duties. The person able to configure an AI system (prompts, thresholds, retrieval sources) is not the person who approves its output for Tier 1 and Tier 2 uses.

6.4 Change control. Prompts, thresholds, model versions, and retrieval sources are governed configurations with version history, documented approval, and rollback plans. Vendor-driven model updates trigger revalidation before continued Tier 1 reliance, since — as COSO notes — model behavior can change without visible notice.

6.5 Evidence and traceability. For any output that lands in the financials or in board materials, [Company] retains the prompt, the output, the model and configuration version, and the identity of who ran and who reviewed it. The standing test: if the auditor asked tomorrow how a number was produced, the answer takes one page.

6.6 Guardrails for wrong answers. Tier 1 and Tier 2 uses define materiality thresholds, escalation paths for outputs that look off, and a fallback so the close still runs the month a tool goes down. Hallucination guardrails — source citations, confidence thresholds, blocked actions below acceptable confidence — follow COSO's control-activity guidance.

7. AI and financial reporting (ICFR)

COSO's 2026 guidance adopts a working definition under which reliance exists when management depends on AI-generated output as part of the evidence supporting an ICFR control's design or operating effectiveness. [Company] adopts the same definition.

Where reliance exists, the AI control pattern must meet the same evidence standards as any ICFR control, including documented prompt, configuration and model version, a clear sampling rationale with exception resolution, and retention of evidence sufficient to support design and operation. Where a control owner independently re-performs the work, evidence expectations may be proportionately lower. The determination of reliance versus non-reliance is made per control by [Controller/CFO], documented in [system], and shared with the external auditor.

8. Third-party and vendor requirements

Any GenAI capability obtained from a vendor limits visibility into training data, update cadence, and underlying controls — a risk COSO flags directly. Before Tier 1 or Tier 2 use, vendors must provide [SOC report / security attestation / contractual notification obligations for model changes / data-use and retention terms], reviewed by [owner]. Vendor model updates are treated as changes under Section 6.4.

9. Monitoring, metrics, and deficiencies

Tier 1 and Tier 2 uses are monitored continuously against defined indicators — accuracy, exception volumes, drift — with separate periodic evaluations (model effectiveness reviews, challenge sessions) to catch what continuous monitoring misses, consistent with COSO Principles 16 and 17. Deficiencies are evaluated for severity, root cause, and scope; deficiencies with material impact are communicated to the audit committee with an action plan and timeline. Root-cause categories include the GenAI-specific failure modes COSO identifies: configuration errors, retrieval issues, prompt design, data quality, model changes outside change control, and vendor changes.

10. Incident response and escalation

AI-related incidents — data leakage, a material wrong output relied upon, prompt-injection exploitation, unauthorized deployment — follow [Company]'s incident response plan with an AI-specific runbook maintained by [owner]. Incidents touching financial reporting or disclosure escalate to [CFO] and the audit committee chair within [timeframe].

11. Training and competence

All personnel receive role-appropriate GenAI training: acceptable-use boundaries and escalation for most staff; secure prompting, bias mitigation, and change control for builders; interpretation of model performance for reviewers and managers. Training recurs [annually] and after material policy changes. Reviewers of Tier 1 output must demonstrate they can evaluate the work without the tool.

12. Policy review

This policy is reviewed [semi-annually] by the AI governance committee and re-approved by the audit committee [annually], or sooner upon material change in technology, regulation, or [Company]'s use of AI. COSO's guidance is explicit that GenAI governance must be revisited faster than traditional policy cycles because the underlying systems change faster.

PART B

Risk-Appetite Statement (template)

Drafted to be read aloud in a committee meeting. Adapt the tiers to the classification in Section 5. COSO's May 2020 guidance, Risk Appetite — Critical to Success, is the reference framework.

Overall statement. [Company] pursues the productivity and insight benefits of generative AI, and accepts the measured risks of adoption, within one absolute boundary: management must at all times be able to explain, evidence, and stand behind every number and statement in its financial reporting and public disclosures, regardless of what tools produced the underlying work.

No appetite. [Company] has no appetite for: AI output entering the general ledger, a filing, or a board document without a named human approval; use of prohibited data categories in unapproved tools; or AI systems making decisions the law reserves for humans.

Low appetite. [Company] has low appetite for reliance on AI within ICFR-relevant processes, and accepts such reliance only where the full evidence standard in Section 7 is met and the external auditor has been informed.

Moderate appetite. [Company] accepts moderate risk in operational (Tier 2) uses where controls in Section 6 operate, exceptions route to humans, and errors are recoverable before they compound.

Higher appetite. [Company] encourages experimentation in productivity and drafting (Tier 3) uses inside acceptable-use boundaries, because — per COSO's observation that GenAI has a low barrier to entry — the realistic alternative to sanctioned experimentation is shadow AI.

Risk responses follow the categories in COSO's ERM framework — accept, avoid, pursue, reduce, share — selected per use case by the AI governance committee and documented in the inventory. Risk accepted outside this appetite requires audit committee approval before, not after, deployment.

PART C

Before the committee signs

A policy signed without questions is a policy nobody read. These are the questions a well-run audit committee should put to management before signing, adapted from the oversight considerations in both COSO publications:

- Do we have a complete inventory of AI use — including shadow AI — and who owns it?
- Which of our AI uses touch financial reporting today, and for each, are we relying on the AI as control evidence or re-performing the work?
- Who is the single executive accountable for AI risk, and what reporting do we receive from them?
- What happens when a vendor updates a model we rely on — do we find out before or after it changes our numbers?
- Has our external auditor seen this policy, and does our evidence retention match what they will ask for?
- What incident involving AI would reach this committee within [timeframe], and has that path been tested?

When management can answer all six without reaching for a binder, the policy is ready to sign.

Approval

Reviewed and approved by the Audit Committee of [Company]:

_____ Chair, Audit Committee Date: _____

_____ Chief Financial Officer Date: _____

_____ [Executive accountable for AI risk] Date: _____

Sources

- COSO, *Achieving Effective Internal Control Over Generative AI (GenAI)*, 2026.
- COSO, *Realize the Full Potential of Artificial Intelligence: Applying the COSO Framework and Principles to Help Implement and Scale Artificial Intelligence*, September 2021.
- COSO, *Risk Appetite — Critical to Success: Using Risk Appetite to Thrive in a Changing World*, May 2020 (referenced within the 2021 publication).

What this template doesn't do

A template gets you a document. It doesn't tier your actual use cases, design the controls behind Section 6, set thresholds your auditor will accept, or build the bench of people who can review AI output without leaning on the tool. That work is judgment work, and it's specific to your company.

If your company is heading toward an audit, a raise, or a sale and this document raised questions you can't yet answer, that conversation is what I do.

Chantal Schutz, CPA, CA · Fractional CFO · chantal.schutz@me.com · chantalschutz.com