

BUILDING TRUST, MANAGING RISK, AND ENABLING SCALE IN ENTERPRISE AI ADOPTION

Responsible AI Governance Framework for Enterprises

Executive briefing — C-suite, board, and compliance leaders | ~20 min

Why Governance Matters Now

- AI is moving from pilot projects to decisions that affect customers, employees, and capital
- Regulators, auditors, and courts are now testing AI outcomes, not just intentions
- Ungoverned deployment creates legal, reputational, and financial exposure at board level
- Competitors with mature governance will out-scale those without it — speed depends on trust
- This is a governance gap, not a technology gap, and it is closing on a fixed regulatory clock

Regulatory Landscape Overview

- EU AI Act: risk-tiered obligations, phased enforcement through 2026-2027, extraterritorial reach
- NIST AI Risk Management Framework: voluntary US baseline, increasingly referenced in contracts and audits
- Sector rules already apply: FTC unfair/deceptive practices, EEOC on hiring algorithms, banking model-risk guidance (SR 11-7)
- State-level activity growing (e.g. Colorado AI Act, NYC Local Law 144 on automated hiring tools)
- No single global standard — compliance strategy must map obligations across jurisdictions, not chase one law

Internal AI Risk Taxonomy

- Model risk: accuracy, drift, hallucination, and unexplainable outputs
- Data risk: privacy violations, IP contamination, unauthorized training data use
- Fairness risk: disparate impact on protected classes in decisions affecting people
- Operational risk: over-reliance, shadow AI usage, vendor lock-in and concentration
- Reputational and legal risk: outputs that breach brand standards, contracts, or regulation

Model Approval Workflow

- Every model or use case is registered before deployment — no exceptions, including vendor tools
- Risk tier assigned at intake (low/medium/high) based on impact and autonomy of the use case
- High-risk use cases require documented testing, legal sign-off, and ethics board review
- Approval is scoped to a specific use case and data context, not a blanket model clearance
- Re-approval triggered by material changes: model version, data source, or use-case expansion

Bias and Fairness Testing Process

- Pre-deployment testing against protected-class subgroups using representative held-out data
- Standard metrics applied: demographic parity, equal opportunity, disparate impact ratio
- Testing repeated at defined intervals post-deployment, not just at launch
- Findings routed to model owners with remediation deadlines and escalation paths
- Independent review for high-stakes use cases (hiring, credit, healthcare) before go-live

Data Privacy Controls

- Data minimization: only collect and process what the use case explicitly requires
- Purpose limitation enforced technically — training and inference data scoped by approved use
- PII handling aligned to applicable law (GDPR, CCPA/CPRA, sector rules) with documented lawful basis
- Vendor and third-party model contracts prohibit using enterprise data for external model training
- Data lineage tracked end-to-end to support audit, deletion, and breach-response obligations

Roles: AI Council and Ethics Board

- AI Council: cross-functional body (legal, risk, IT, business units) owns policy and approval standards
- Ethics Board: independent review of high-risk or ethically ambiguous use cases, escalation authority
- Executive sponsor accountable to the board for program outcomes and material incidents
- Business unit owners accountable for day-to-day compliance of their deployed models
- Clear escalation path from model owner through council to ethics board to executive sponsor

Vendor and Third-Party AI Risk Assessment

- Standard due-diligence questionnaire covering training data provenance, security, and bias testing
- Contractual requirements: audit rights, incident notification, data-use restrictions, indemnification
- Vendor risk tier assigned based on data sensitivity and decision impact, mirroring internal taxonomy
- Ongoing monitoring, not one-time diligence — vendor model changes trigger re-assessment
- Concentration risk tracked across the vendor portfolio, not just per-vendor

Audit and Monitoring Approach

- Continuous monitoring of production models for drift, degraded accuracy, and anomalous outputs
- Independent internal audit function reviews governance program effectiveness annually
- Audit trail maintained for every model decision: inputs, version, approval record, override history
- Monitoring thresholds trigger automatic alerts to model owners and risk function
- External audit or attestation considered for highest-risk, regulated use cases

Incident Response Protocol

- Defined severity classification for AI incidents, from minor output error to regulatory-reportable event
- Cross-functional response team activated within a fixed time window of detection
- Immediate containment options: model rollback, feature disablement, human-in-the-loop override
- Root cause analysis and remediation plan required before re-enabling affected use case
- Regulatory and customer notification procedures pre-defined, not improvised under pressure

Training and Awareness Program

- Mandatory baseline training for all employees using AI tools, role-specific modules beyond that
- Deep-dive training for model owners, developers, and legal/compliance on governance obligations
- Board and executive briefings on AI risk posture, at least annually or on material change
- Practical guidance on acceptable use, shadow AI risks, and escalation channels
- Training effectiveness tracked and refreshed as regulation and internal policy evolve

Case Study: A Governance Failure Elsewhere

- A widely reported case: a hiring algorithm found to systematically penalize resumes mentioning women's colleges
- Root cause was training data reflecting historical hiring bias, not intentional design
- The tool was used for years before the pattern was identified and the project was discontinued
- Key lesson: absence of pre-deployment fairness testing let a biased pattern scale undetected
- Governance controls in this deck — testing, review, monitoring — map directly to preventing this failure mode

Maturity Roadmap

- Stage 1 – Foundational: inventory of AI use cases, baseline policy, risk taxonomy adopted
- Stage 2 – Managed: approval workflow live, fairness testing standardized, council operating
- Stage 3 – Integrated: monitoring and audit embedded, vendor risk program mature
- Stage 4 – Optimized: governance data informs strategy, program benchmarked externally
- Target state and timeline should be set against current stage, not assumed to be Stage 4

Call to Action

- Approve the AI Council charter and ethics board mandate this quarter
- Fund the model registry and approval workflow as the first operational priority
- Mandate baseline AI training completion for all employees within 90 days
- Commission an inventory of all current AI use cases, including shadow deployments
- Schedule quarterly board-level review of AI risk posture and incident metrics