



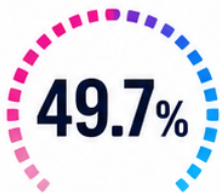
MODERN AI
smart. secure. transformative

AI GOVERNANCE & RESPONSIBLE AI

Executive Assessment Report



Sample Client Current-State Assessment
Based on 72 controls across 12 governance
and responsible AI domains



Overall residual risk



Risk rating



72

Controls assessed



60

Critical / high findings



Prepared by



MODERN AI
smart. secure. transformative



Table of Contents

This report uses Word heading styles. In Microsoft Word, select References > Table of Contents to insert or refresh an automatic table of contents.

1. Executive Summary
2. Assessment Approach and Limitations
3. Enterprise Risk and Maturity Profile
4. Domain-by-Domain Findings
5. Governance Artifact Readiness
6. AI Risk Register Analysis
7. Target Operating Model and Architecture
8. Prioritized Recommendations
9. 120-Day Improvement Roadmap
10. Conclusion



1. Executive Summary

The assessment evaluates 72 controls across 12 AI governance and responsible AI domains. The populated workbook produces an overall residual risk score of 49.7% and an overall rating of Medium. Sixty controls are rated Critical or High, indicating that the organization has several useful foundational practices but lacks the policies, evidence, repeatability, assurance, and operating discipline required for a mature enterprise AI governance program.

The strongest area is Data Governance, Privacy & Consent at 7.4% residual risk. The weakest areas are Regulatory Compliance & Audit Evidence at 80.0% and Responsible AI Principles & Policy at 75.0%. Monitoring, incident response, training, model governance, fairness, human oversight, and AI security also require accelerated remediation. All twelve governance artifacts are marked Not Created, and all eight enterprise AI risks are Not Started, confirming a significant implementation gap between current practices and an auditable target state.

Executive Conclusions

- AI governance exists as a developing set of activities rather than a consistently governed enterprise operating model.
- Policy and regulatory evidence gaps create the greatest immediate legal, audit, reputational, and executive-accountability exposure.
- Technical and lifecycle controls are only partially implemented across model risk, fairness, oversight, security, and production monitoring.
- The program requires an authoritative AI inventory, risk-tiered intake process, named owners, approval gates, evidence retention, and board-level reporting.
- The workbook already provides a practical 120-day remediation sequence; execution now depends on assigned owners, target dates, funding, and operating cadence.

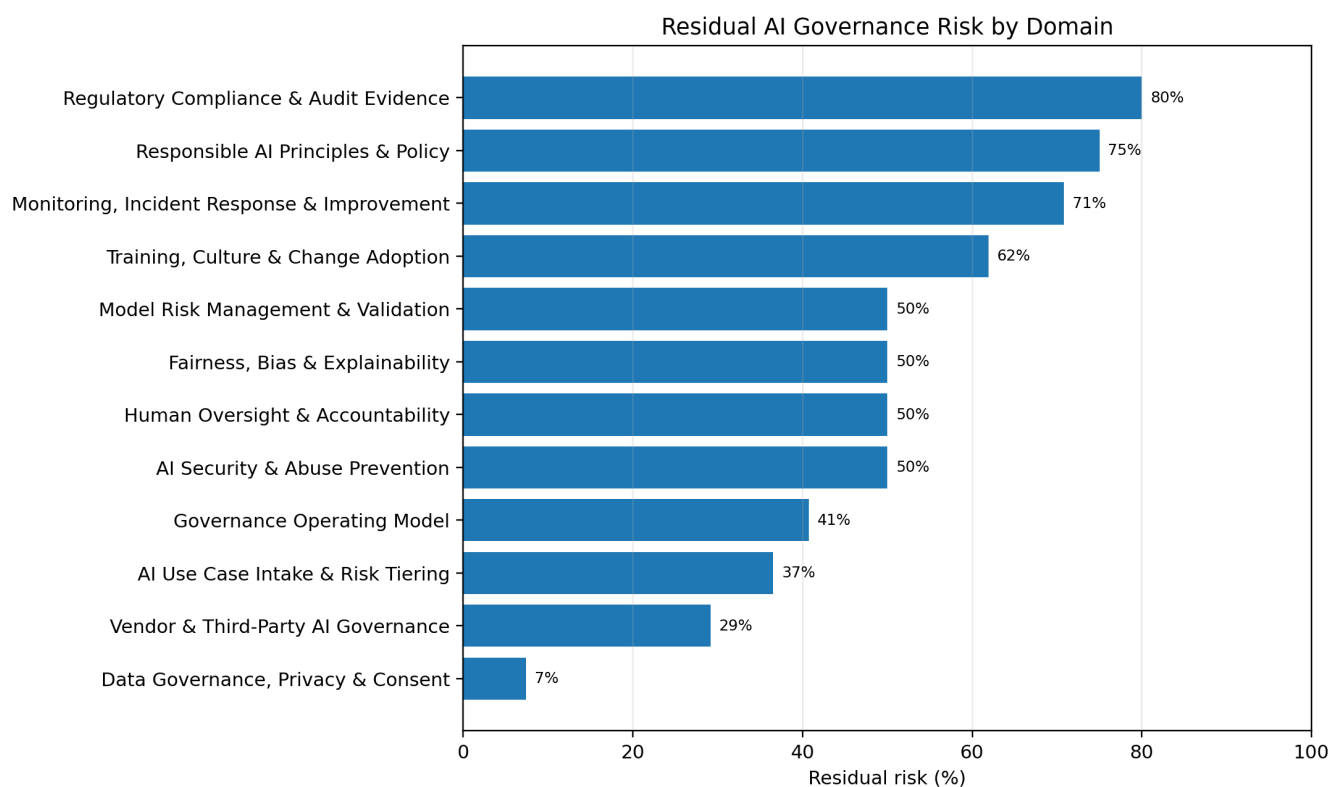


Figure 1. Residual risk by AI governance domain.



2. Assessment Approach and Limitations

The workbook uses Yes, Partial, No, and Unknown responses to calculate control scores and weighted residual risk. Each control also includes priority, risk rating, evidence strength, maturity level, ownership, target date, status, framework alignment, and recommended evidence. Domain risk is the sum of weighted residual risk divided by the maximum possible weighted risk for that domain.

Response	Control score	Interpretation
Yes	1.0	Control appears implemented; validate evidence and operating effectiveness.
Partial	0.5	Control exists in part but is inconsistent, incomplete, or weakly evidenced.
No / Unknown	0.0	Control is absent or cannot be demonstrated; treat as a remediation requirement.

Assessment Limitations

- Client name, assessment date, owners, target dates, consultant notes, and evidence locations are not populated.
- All controls show a Developing maturity label and Not Started status, so the assessment should be treated as a sample baseline rather than a formal assurance opinion.
- A Yes response should not be considered fully effective until evidence is reviewed and the control is tested in operation.
- Risk ratings reflect the workbook methodology and should be calibrated to the client's legal obligations, risk appetite, AI use cases, data sensitivity, and autonomy levels.

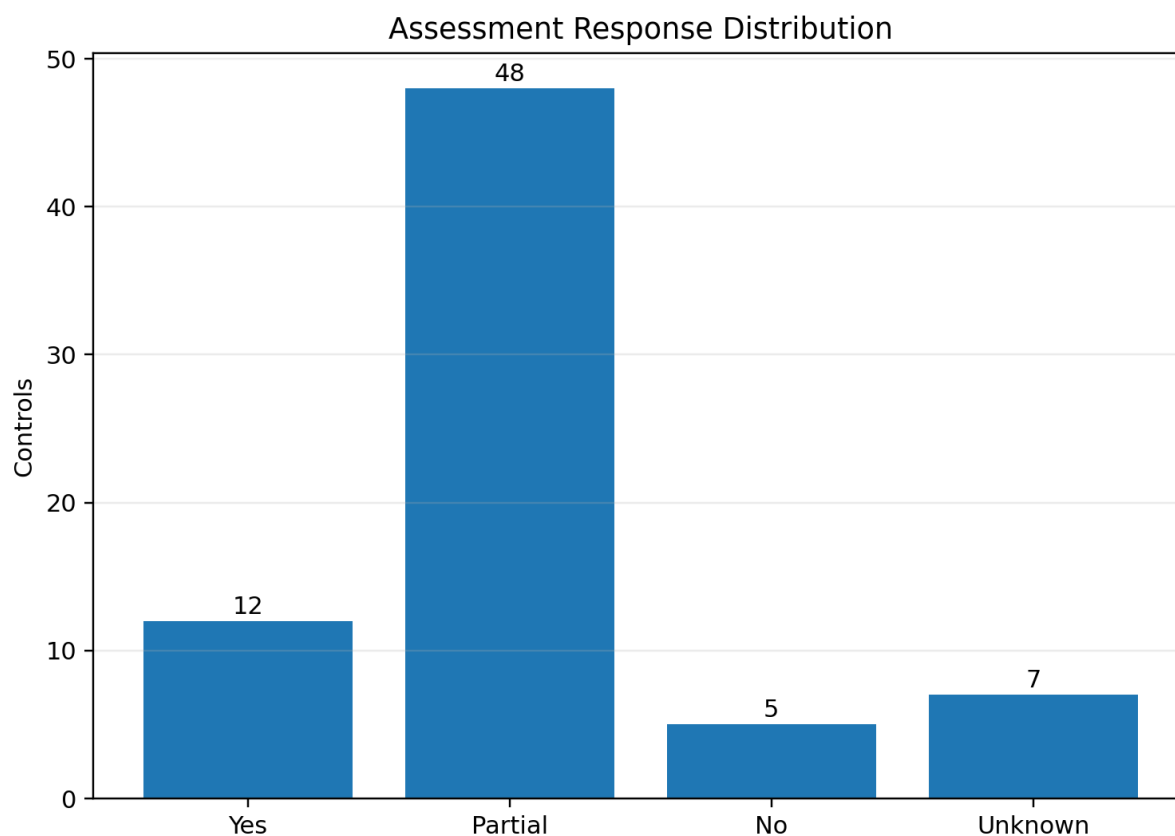


Figure 2. Distribution of assessment responses.



3. Enterprise Risk and Maturity Profile

2	6	3	1
Critical domains	High domains	Medium domains	Low domains

3.1 Risk Concentration

The risk profile is concentrated in policy, auditability, and operational assurance. These domains influence all AI use cases, so weakness can propagate across the entire portfolio. A critical policy or evidence gap cannot be solved by technology alone; it requires executive authority, approved standards, repeatable workflows, retained records, and independent testing.

Control Risk Rating Distribution

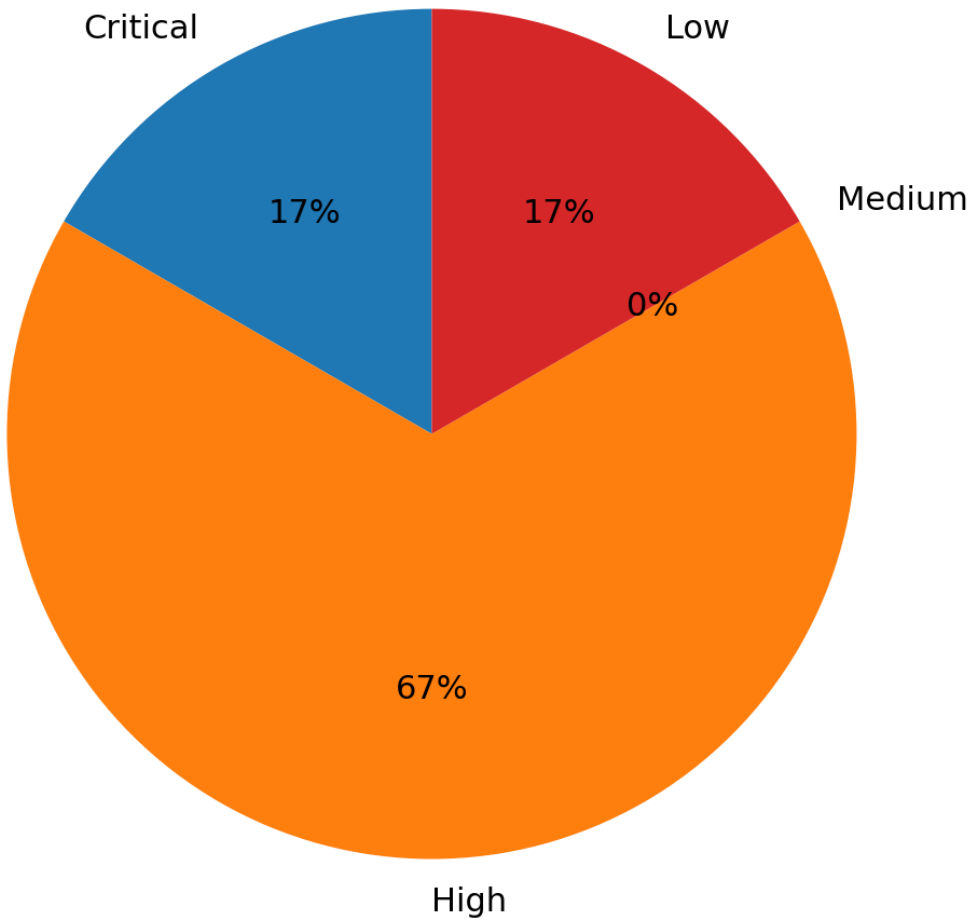


Figure 3. Control-level risk rating distribution.



3.2 Priority Findings

- Regulatory Compliance & Audit Evidence: 80.0% residual risk (Critical). Immediate executive action and control design required.
- Responsible AI Principles & Policy: 75.0% residual risk (Critical). Immediate executive action and control design required.
- Monitoring, Incident Response & Improvement: 70.8% residual risk (High). Prioritize remediation in next 30-60 days.
- Training, Culture & Change Adoption: 61.9% residual risk (High). Prioritize remediation in next 30-60 days.
- Model Risk Management & Validation: 50.0% residual risk (High). Prioritize remediation in next 30-60 days.
- Fairness, Bias & Explainability: 50.0% residual risk (High). Prioritize remediation in next 30-60 days.

3.3 Strengths to Preserve

- RAI-001 - Executive sponsorship: Is there an accountable executive sponsor for enterprise AI governance with decision authority?
- RAI-017 - Prohibited use cases: Has the organization defined prohibited AI use cases and unacceptable risk scenarios?
- RAI-019 - Data classification: Are data classification rules applied to AI inputs, prompts, outputs, embeddings, training data, and logs?
- RAI-020 - Privacy assessment: Are privacy impact assessments completed for AI systems using personal, confidential, or regulated data?
- RAI-021 - Consent and purpose: Is there evidence that personal data used by AI is aligned to consent, purpose limitation, and retention requirements?
- RAI-022 - Sensitive data controls: Are DLP, masking, tokenization, encryption, or redaction controls used to prevent sensitive data leakage into AI systems?
- RAI-023 - Data lineage: Can the organization trace data sources used by AI systems, models, RAG repositories, and decision workflows?
- RAI-052 - Embedded AI review: Are new AI features in existing SaaS platforms reviewed before activation?



4. Domain-by-Domain Findings

4.1 Governance Operating Model

This domain measures executive sponsorship, committee authority, roles, accountability, inventory, and integration with enterprise risk.

6	40.7%	Medium	0.58
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Formalize the AI governance operating model through an approved charter, accountable executive sponsor, cross-functional steering committee, RACI, decision rights, escalation thresholds, and an authoritative AI inventory.

Highest-Priority Control Gaps

- RAI-002 - AI governance committee: Partial response; weighted risk 2.5; High risk. Has the organization established a cross-functional AI governance committee including business, technology, security, legal, privacy, risk, and compliance?
- RAI-005 - AI portfolio visibility: Partial response; weighted risk 2.5; High risk. Does leadership have visibility into all active, planned, and shadow AI initiatives?
- RAI-003 - RACI and decision rights: Partial response; weighted risk 2; High risk. Are AI governance roles, responsibilities, approval authorities, and escalation paths documented in a RACI model?
- RAI-004 - Governance cadence: Partial response; weighted risk 2; High risk. Is there a defined meeting cadence for AI risk review, control approvals, exception management, and portfolio oversight?
- RAI-006 - Board reporting: Partial response; weighted risk 2; High risk. Are AI risk, adoption, incidents, and compliance metrics reported to senior leadership or the board?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-002	AI governance committee	Partial	2.5	Critical
RAI-005	AI portfolio visibility	Partial	2.5	Critical
RAI-003	RACI and decision rights	Partial	2	High
RAI-004	Governance cadence	Partial	2	High
RAI-006	Board reporting	Partial	2	High
RAI-001	Executive sponsorship	Yes	0	Critical



4.2 Responsible AI Principles & Policy

This domain evaluates whether responsible AI principles and acceptable-use expectations are approved, communicated, and translated into enforceable standards.

6	75.0%	Critical	0.25
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

This is a critical weakness. Publish an enterprise Responsible AI Policy and Generative AI Acceptable Use Standard, define prohibited uses and exceptions, and connect principles to measurable controls and employee obligations.

Highest-Priority Control Gaps

- RAI-011 - Responsible AI review: No response; weighted risk 5; Critical risk. Are high-impact or sensitive AI use cases reviewed against responsible AI principles before production deployment?
- RAI-009 - Generative AI policy: No response; weighted risk 4; Critical risk. Does the policy specifically address generative AI, copilots, prompts, outputs, hallucinations, and prohibited use cases?
- RAI-010 - Policy exceptions: No response; weighted risk 4; Critical risk. Is there a documented process to request, approve, monitor, and expire exceptions to AI policies?
- RAI-007 - Responsible AI principles: Partial response; weighted risk 2.5; High risk. Has the organization formally defined responsible AI principles such as fairness, transparency, accountability, privacy, safety, and human oversight?
- RAI-008 - AI acceptable use: Partial response; weighted risk 2.5; High risk. Is there an approved enterprise AI acceptable use policy for employees, contractors, and business partners?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-011	Responsible AI review	No	5	Critical
RAI-009	Generative AI policy	No	4	High
RAI-010	Policy exceptions	No	4	High
RAI-007	Responsible AI principles	Partial	2.5	Critical
RAI-008	AI acceptable use	Partial	2.5	Critical
RAI-012	Policy communication	Partial	1.5	Medium



4.3 AI Use Case Intake & Risk Tiering

This domain assesses whether AI initiatives enter through a consistent intake, classification, approval, and risk-tiering workflow.

6	36.5%	Medium	0.67
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Implement a single intake workflow that records business value, data, users, autonomy, impacts, vendor dependencies, risk tier, approval route, and required control evidence before development or activation.

Highest-Priority Control Gaps

- RAI-013 - Use case intake: Partial response; weighted risk 2.5; High risk. Is there a formal intake process for new AI use cases, including business objective, data used, users impacted, and expected outcomes?
- RAI-014 - Risk tiering model: Partial response; weighted risk 2.5; High risk. Does the organization classify AI use cases by inherent risk level, such as low, medium, high, or prohibited?
- RAI-015 - Impact assessment: Partial response; weighted risk 2.5; High risk. Are AI impact assessments performed for use cases that affect customers, employees, credit, hiring, health, safety, or regulated decisions?
- RAI-016 - Approval gates: Partial response; weighted risk 2; High risk. Are there defined approval gates before pilot, production, external release, or autonomous decision-making?
- RAI-017 - Prohibited use cases: Yes response; weighted risk 0; Low risk. Has the organization defined prohibited AI use cases and unacceptable risk scenarios?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-013	Use case intake	Partial	2.5	Critical
RAI-014	Risk tiering model	Partial	2.5	Critical
RAI-015	Impact assessment	Partial	2.5	Critical
RAI-016	Approval gates	Partial	2	High
RAI-017	Prohibited use cases	Yes	0	High
RAI-018	Value tracking	Yes	0	Medium



4.4 Data Governance, Privacy & Consent

This domain assesses data classification, lawful use, consent, retention, lineage, quality, and controls over prompts, outputs, and training data.

6	7.4%	Low	0.92
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

This is the strongest domain in the sample. Maintain effectiveness through DLP, masking, approved data boundaries, retention rules, privacy impact assessments, and periodic testing of prompt and output handling.

Highest-Priority Control Gaps

- RAI-024 - Retention and deletion: Partial response; weighted risk 2; High risk. Are retention and deletion rules defined for prompts, outputs, model artifacts, evaluation data, and AI logs?
- RAI-019 - Data classification: Yes response; weighted risk 0; Low risk. Are data classification rules applied to AI inputs, prompts, outputs, embeddings, training data, and logs?
- RAI-020 - Privacy assessment: Yes response; weighted risk 0; Low risk. Are privacy impact assessments completed for AI systems using personal, confidential, or regulated data?
- RAI-021 - Consent and purpose: Yes response; weighted risk 0; Low risk. Is there evidence that personal data used by AI is aligned to consent, purpose limitation, and retention requirements?
- RAI-022 - Sensitive data controls: Yes response; weighted risk 0; Low risk. Are DLP, masking, tokenization, encryption, or redaction controls used to prevent sensitive data leakage into AI systems?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-024	Retention and deletion	Partial	2	High
RAI-019	Data classification	Yes	0	Critical
RAI-020	Privacy assessment	Yes	0	Critical
RAI-021	Consent and purpose	Yes	0	High
RAI-022	Sensitive data controls	Yes	0	Critical
RAI-023	Data lineage	Yes	0	High



4.5 Model Risk Management & Validation

This domain covers inventory, documentation, validation, performance thresholds, change control, monitoring, and retirement.

6	50.0%	High	0.50
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Create a model risk standard and model/system cards, independent validation criteria, release gates, change controls for models/prompts/RAG content, performance thresholds, and controlled retirement.

Highest-Priority Control Gaps

- RAI-025 - Model inventory: Partial response; weighted risk 2.5; High risk. Is there a centralized inventory of AI models, foundation models, fine-tuned models, APIs, copilots, and automated decision systems?
- RAI-027 - Validation process: Partial response; weighted risk 2.5; High risk. Are models independently validated before production based on accuracy, robustness, fairness, privacy, security, and intended use?
- RAI-026 - Model documentation: Partial response; weighted risk 2; High risk. Are model cards, system cards, or equivalent documentation maintained for material AI systems?
- RAI-028 - Change control: Partial response; weighted risk 2; High risk. Are model updates, prompt changes, RAG content changes, and configuration changes managed through change control?
- RAI-029 - Performance thresholds: Partial response; weighted risk 2; High risk. Are acceptance criteria and performance thresholds defined for model quality, error tolerance, drift, and safety?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-025	Model inventory	Partial	2.5	Critical
RAI-027	Validation process	Partial	2.5	Critical
RAI-026	Model documentation	Partial	2	High
RAI-028	Change control	Partial	2	High
RAI-029	Performance thresholds	Partial	2	High
RAI-030	Retirement process	Partial	1.5	Medium



4.6 Fairness, Bias & Explainability

This domain evaluates bias testing, fairness measures, transparency, explainability, disclosure, appeals, and stakeholder harm.

6	50.0%	High	0.50
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Define when fairness and explainability controls are mandatory, select relevant metrics, perform impact assessments, provide user disclosure and contestability, and document residual harm and executive acceptance.

Highest-Priority Control Gaps

- RAI-031 - Bias testing: Partial response; weighted risk 2.5; High risk. Are AI systems tested for bias or discriminatory outcomes where decisions impact people, access, pricing, service, employment, or eligibility?
- RAI-035 - Appeals process: Partial response; weighted risk 2.5; High risk. Is there a mechanism for individuals or employees to challenge, appeal, or escalate AI-assisted decisions?
- RAI-032 - Fairness metrics: Partial response; weighted risk 2; High risk. Are appropriate fairness metrics selected, documented, and monitored for high-risk AI systems?
- RAI-033 - Explainability requirements: Partial response; weighted risk 2; High risk. Are explainability requirements defined for AI systems that influence material business or customer outcomes?
- RAI-034 - User disclosures: Partial response; weighted risk 2; High risk. Are users informed when they are interacting with AI or when AI materially informs a decision?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-031	Bias testing	Partial	2.5	Critical
RAI-035	Appeals process	Partial	2.5	Critical
RAI-032	Fairness metrics	Partial	2	High
RAI-033	Explainability requirements	Partial	2	High
RAI-034	User disclosures	Partial	2	High
RAI-036	Harm assessment	Partial	2	High



4.7 Human Oversight & Accountability

This domain assesses human-in-the-loop requirements, ownership, decision boundaries, override capability, accountability logs, and operating playbooks.

6	50.0%	High	0.50
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Assign named business and technical owners to every AI system, define limits on autonomous action, require human review for material decisions, and preserve approvals, overrides, exceptions, and agent actions.

Highest-Priority Control Gaps

- RAI-037 - Human-in-the-loop: Partial response; weighted risk 2.5; High risk. Are human review and approval requirements defined for high-risk, high-impact, or autonomous AI decisions?
- RAI-038 - Accountable owner: Partial response; weighted risk 2.5; High risk. Does each AI system have a named business owner and technical owner accountable for outcomes and controls?
- RAI-039 - Decision authority: Partial response; weighted risk 2.5; High risk. Are limits defined for what AI can recommend, decide, automate, or execute without human approval?
- RAI-040 - Override procedures: Partial response; weighted risk 2; High risk. Can authorized humans override, pause, disable, or roll back AI-assisted workflows?
- RAI-041 - Accountability logging: Partial response; weighted risk 2; High risk. Are approvals, overrides, agent actions, exceptions, and high-risk decisions logged for audit?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-037	Human-in-the-loop	Partial	2.5	Critical
RAI-038	Accountable owner	Partial	2.5	Critical
RAI-039	Decision authority	Partial	2.5	Critical
RAI-040	Override procedures	Partial	2	High
RAI-041	Accountability logging	Partial	2	High
RAI-042	Operational playbooks	Partial	1.5	Medium



4.8 AI Security & Abuse Prevention

This domain covers AI threat modeling, secure development, identity, prompt/output controls, secrets protection, monitoring, and abuse prevention.

6	50.0%	High	0.50
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Integrate AI security into the secure SDLC using threat modeling, least privilege, MFA/RBAC/PAM, secrets management, prompt and output controls, red teaming, logging, SIEM integration, and AI-specific incident playbooks.

Highest-Priority Control Gaps

- RAI-043 - AI threat model: Partial response; weighted risk 2.5; High risk. Are AI systems threat modeled for prompt injection, data leakage, model abuse, malicious outputs, agent misuse, and supply chain compromise?
- RAI-045 - Access controls: Partial response; weighted risk 2.5; High risk. Are AI platforms protected using strong identity, MFA, RBAC, PAM, service account governance, and least privilege?
- RAI-047 - Secrets protection: Partial response; weighted risk 2.5; High risk. Are API keys, model credentials, connectors, tokens, and agent secrets managed securely?
- RAI-044 - Secure SDLC: Partial response; weighted risk 2; High risk. Are AI applications, models, prompts, and agent workflows integrated into secure development and release practices?
- RAI-046 - Prompt/output controls: Partial response; weighted risk 2; High risk. Are prompt filtering, output validation, content safety, jailbreak resistance, and sensitive data blocking implemented where required?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-043	AI threat model	Partial	2.5	Critical
RAI-045	Access controls	Partial	2.5	Critical
RAI-047	Secrets protection	Partial	2.5	Critical
RAI-044	Secure SDLC	Partial	2	High
RAI-046	Prompt/output controls	Partial	2	High
RAI-048	Security monitoring	Partial	2	High



4.9 Vendor & Third-Party AI Governance

This domain evaluates third-party inventory, due diligence, contract terms, embedded SaaS AI review, assurance evidence, and exit planning.

6	29.2%	Medium	0.75
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Implement AI-specific vendor due diligence and contract clauses covering data use, model training, residency, security, audit rights, sub-processors, incidents, portability, and exit.

Highest-Priority Control Gaps

- RAI-049 - Vendor inventory: Partial response; weighted risk 2.5; High risk. Is there an inventory of third-party AI tools, embedded SaaS AI features, model providers, and AI subcontractors?
- RAI-050 - Vendor risk assessment: Partial response; weighted risk 2.5; High risk. Are AI vendors assessed for data use, model training terms, privacy, security, residency, audit rights, support access, and subprocessors?
- RAI-051 - Contract clauses: Partial response; weighted risk 2; High risk. Do contracts include AI-specific clauses covering data ownership, model training, retention, explainability, incident notification, and exit rights?
- RAI-052 - Embedded AI review: Yes response; weighted risk 0; Low risk. Are new AI features in existing SaaS platforms reviewed before activation?
- RAI-053 - Third-party evidence: Yes response; weighted risk 0; Low risk. Does the organization collect vendor evidence such as SOC 2, ISO 27001, ISO 42001, model documentation, security testing, and data flow diagrams?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-049	Vendor inventory	Partial	2.5	Critical
RAI-050	Vendor risk assessment	Partial	2.5	Critical
RAI-051	Contract clauses	Partial	2	High
RAI-052	Embedded AI review	Yes	0	High
RAI-053	Third-party evidence	Yes	0	Medium
RAI-054	Exit strategy	Yes	0	Medium



4.10 Regulatory Compliance & Audit Evidence

This domain measures obligations mapping, framework alignment, legal review, evidence retention, control testing, and records management.

6	80.0%	Critical	0.17
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

This is the highest-risk domain. Build a regulatory obligations and control crosswalk, establish legal review triggers, define evidence standards and retention, and begin recurring design and operating-effectiveness testing.

Highest-Priority Control Gaps

- RAI-057 - Audit evidence: Unknown response; weighted risk 5; Critical risk. Is evidence collected and retained for AI approvals, testing, monitoring, incidents, exceptions, and training?
- RAI-056 - Framework alignment: Unknown response; weighted risk 4; Critical risk. Is the AI governance program aligned to frameworks such as NIST AI RMF, ISO/IEC 42001, ISO 27001, SOC 2, OSFI B-13, or applicable AI laws?
- RAI-058 - Control testing: Unknown response; weighted risk 4; Critical risk. Are AI governance controls periodically tested for design and operating effectiveness?
- RAI-059 - Legal review: Unknown response; weighted risk 4; Critical risk. Does legal review AI use cases for regulatory, contractual, intellectual property, privacy, and liability implications?
- RAI-060 - Records management: Unknown response; weighted risk 3; Critical risk. Are AI governance records managed according to retention, eDiscovery, audit, and records requirements?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-057	Audit evidence	Unknown	5	Critical
RAI-056	Framework alignment	Unknown	4	High
RAI-058	Control testing	Unknown	4	High
RAI-059	Legal review	Unknown	4	High
RAI-060	Records management	Unknown	3	Medium
RAI-055	Regulatory mapping	Yes	0	Critical



4.11 Monitoring, Incident Response & Improvement

This domain assesses production monitoring, AI incidents, KRIs/KPIs, feedback, post-implementation review, and continual improvement.

6	70.8%	High	0.33
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Deploy monitoring for drift, hallucinations, harmful outputs, policy violations, security events, overrides, and performance. Establish AI incident scenarios, thresholds, escalation, lessons learned, and post-launch review.

Highest-Priority Control Gaps

- RAI-061 - AI monitoring: Unknown response; weighted risk 5; Critical risk. Are AI systems monitored for drift, degradation, hallucination rates, policy violations, bias, security events, and abnormal behavior?
- RAI-062 - Incident playbooks: Unknown response; weighted risk 5; Critical risk. Are there AI-specific incident response playbooks for data leakage, harmful outputs, rogue agents, model compromise, and unauthorized AI use?
- RAI-063 - KRIs and KPIs: Partial response; weighted risk 2; High risk. Are AI risk indicators and performance indicators tracked for governance reporting?
- RAI-065 - Continuous improvement: Partial response; weighted risk 2; High risk. Is there a continuous improvement process to update governance as laws, technologies, and risks evolve?
- RAI-064 - Feedback loop: Partial response; weighted risk 1.5; High risk. Is user feedback captured and used to improve AI systems, policies, and controls?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-061	AI monitoring	Unknown	5	Critical
RAI-062	Incident playbooks	Unknown	5	Critical
RAI-063	KRIs and KPIs	Partial	2	High
RAI-065	Continuous improvement	Partial	2	High
RAI-064	Feedback loop	Partial	1.5	Medium
RAI-066	Post-implementation review	Partial	1.5	Medium



4.12 Training, Culture & Change Adoption

This domain assesses role-based training, employee awareness, developer enablement, leadership sponsorship, change management, and recurring communication.

6	61.9%	High	0.33
Questions	Residual risk	Risk rating	Avg control score

Executive Interpretation

Create a role-based education program for executives, employees, developers, risk teams, security, legal, procurement, and operators. Use communications and workflow integration to reduce shadow AI and improve adoption.

Highest-Priority Control Gaps

- RAI-071 - Change management: No response; weighted risk 3; Critical risk. Is there a change management plan to drive adoption of AI governance workflows and tools?
- RAI-068 - Employee awareness: Partial response; weighted risk 2.5; High risk. Are employees trained on safe AI usage, data handling, prompt hygiene, prohibited use, and reporting concerns?
- RAI-067 - Role-based training: Partial response; weighted risk 2; High risk. Is AI governance training tailored by role for executives, developers, business users, risk, legal, privacy, and security teams?
- RAI-069 - Developer enablement: Partial response; weighted risk 2; High risk. Are developers and AI builders trained on responsible AI design, testing, documentation, and secure AI development?
- RAI-072 - Communication plan: No response; weighted risk 2; Critical risk. Are responsible AI expectations communicated regularly through newsletters, intranet, lunch-and-learns, or policy refreshers?

Recommended Evidence and Management Action

- Assign a named business owner and control owner.
- Document the approved policy, standard, procedure, workflow, or technical configuration.
- Retain evidence in a controlled repository with review and retention requirements.
- Test design and operating effectiveness, track exceptions, and report overdue remediation.

ID	Assessment area	Answer	Weighted risk	Priority
RAI-071	Change management	No	3	Medium
RAI-068	Employee awareness	Partial	2.5	Critical
RAI-067	Role-based training	Partial	2	High
RAI-069	Developer enablement	Partial	2	High
RAI-072	Communication plan	No	2	Low
RAI-070	Leadership adoption	Partial	1.5	Medium



5. Governance Artifact Readiness

The workbook identifies twelve foundational governance artifacts. All are marked Not Created and High priority. This represents a material control-design gap because these artifacts establish the policies, standards, records, and operating expectations needed to make governance repeatable and auditable.

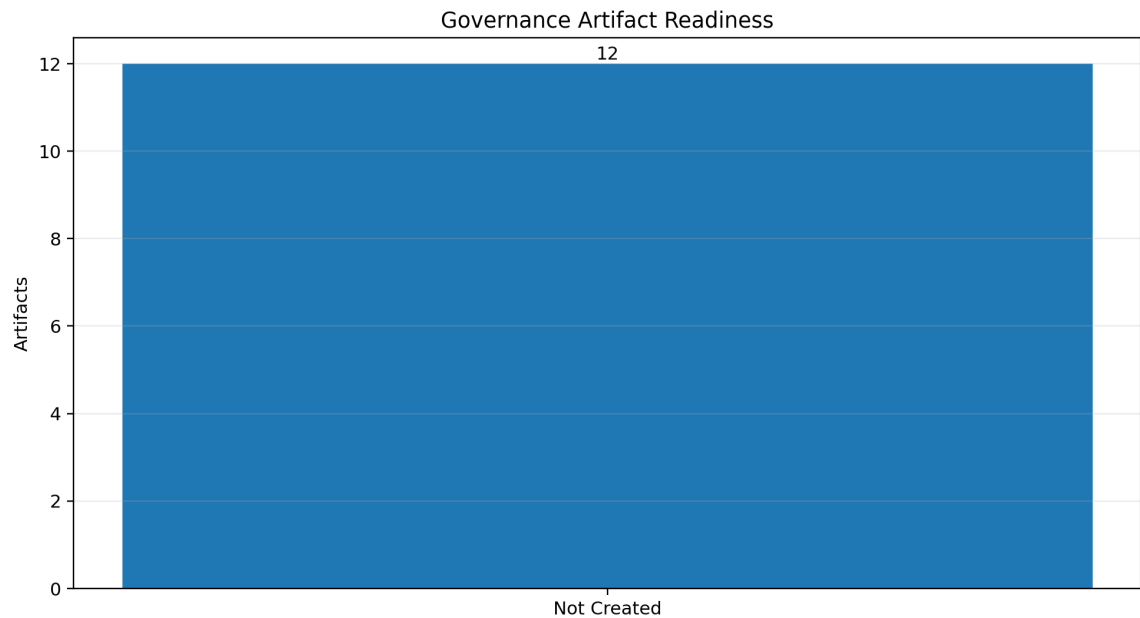


Figure 4. Governance artifact status.

ID	Governance artifact	Purpose	Status	Priority
GOV-001	Enterprise AI Governance Charter	Defines decision rights, accountability, committee scope, and escalation paths.	Not Created	High
GOV-002	Responsible AI Policy	Defines responsible AI principles, employee expectations, and prohibited usage.	Not Created	High
GOV-003	Generative AI Acceptable Use Standard	Specific guardrails for public/private AI tools, prompts, outputs, and sensitive data.	Not Created	High
GOV-004	AI Use Case Intake Form	Captures purpose, business value, data, users, risk tier, and required approvals.	Not Created	High
GOV-005	AI Risk Tiering Methodology	Classifies AI use cases by risk and maps required controls and approvals.	Not Created	High
GOV-006	AI Impact Assessment Template	Assesses privacy, fairness, safety, security, explainability, and stakeholder impact.	Not Created	High
GOV-007	AI Model Risk Management Standard	Defines inventory, validation, testing, monitoring, change control, and retirement requirements.	Not Created	High
GOV-008	AI Vendor Risk Questionnaire	Assesses third-party AI data usage, model training, residency, security, and auditability.	Not Created	High



ID	Governance artifact	Purpose	Status	Priority
GOV-009	AI Human Oversight Standard	Defines when human review, approval, override, and escalation are required.	Not Created	High
GOV-010	AI Incident Response Playbook	Defines response steps for AI leakage, harmful outputs, rogue agents, and model compromise.	Not Created	High
GOV-011	AI Monitoring & Audit Evidence Standard	Defines logs, metrics, evidence retention, and audit reporting requirements.	Not Created	High
GOV-012	AI Training & Awareness Plan	Defines role-based training for employees, executives, developers, and risk teams.	Not Created	High

Recommended Artifact Sequence

- Days 1-30: Governance Charter, Responsible AI Policy, Generative AI Acceptable Use Standard, AI Use Case Intake Form.
- Days 31-60: AI Risk Tiering Methodology, AI Impact Assessment, AI Vendor Risk Questionnaire.
- Days 61-90: Model Risk Standard, Human Oversight Standard, AI Incident Response Playbook.
- Days 91-120: Monitoring and Audit Evidence Standard, Training and Awareness Plan, control-testing records.



6. AI Risk Register Analysis

The risk register contains eight representative enterprise AI risks. Three are rated Critical and five High. Every risk is marked Not Started and lacks an owner or target date. The immediate management requirement is to validate inherent risk, map controls, assign accountable owners, approve treatment plans, and establish residual-risk acceptance authority.

Risk	Theme	Domain	Score	Rating	Treatment
RISK-001	Shadow AI	Governance Operating Model	20	Critical	Create AI inventory, discovery process, and approval workflow.
RISK-002	Sensitive data leakage	Data Governance, Privacy & Consent	20	Critical	Implement AI data policy, DLP, masking, and employee training.
RISK-003	Model bias	Fairness, Bias & Explainability	15	High	Perform impact assessments, fairness testing, and human oversight.
RISK-004	Weak vendor control	Vendor & Third-Party AI Governance	16	High	Update vendor questionnaire and contract clauses.
RISK-005	Unclear accountability	Human Oversight & Accountability	12	High	Create AI RACI and require owner assignment for each AI system.
RISK-006	AI security abuse	AI Security & Abuse Prevention	20	Critical	Integrate AI threat modeling, logging, and incident playbooks.
RISK-007	Audit readiness gaps	Regulatory Compliance & Audit Evidence	12	High	Create evidence standard and control testing cadence.
RISK-008	Unmonitored AI behavior	Monitoring, Incident Response & Improvement	16	High	Define KPIs/KRIs, monitoring, alerts, and incident response.

Most Material Risk Themes

- Shadow AI (Critical, score 20): Unapproved AI tools or embedded SaaS AI features are used without governance review. Recommended treatment: Create AI inventory, discovery process, and approval workflow.
- Sensitive data leakage (Critical, score 20): Confidential or personal data may be entered into public AI tools or retained in prompts/logs. Recommended treatment: Implement AI data policy, DLP, masking, and employee training.
- AI security abuse (Critical, score 20): AI applications or agents may be vulnerable to prompt injection, data exfiltration, or abuse. Recommended treatment: Integrate AI threat modeling, logging, and incident playbooks.
- Weak vendor control (High, score 16): Third-party AI vendors may use data for training or operate with insufficient transparency. Recommended treatment: Update vendor questionnaire and contract clauses.
- Unmonitored AI behavior (High, score 16): AI system drift, hallucinations, or policy violations may not be detected in production. Recommended treatment: Define KPIs/KRIs, monitoring, alerts, and incident response.
- Model bias (High, score 15): AI systems could produce unfair or discriminatory outcomes in high-impact use cases. Recommended treatment: Perform impact assessments, fairness testing, and human oversight.
- Unclear accountability (High, score 12): AI systems may lack business owners and technical owners accountable for outcomes. Recommended treatment: Create AI RACI and require owner assignment for each AI system.
- Audit readiness gaps (High, score 12): AI governance decisions and controls may not have retained evidence for audit. Recommended treatment: Create evidence standard and control testing cadence.



7. Target Operating Model and Architecture

The target operating model should connect business value, governance, legal and regulatory obligations, responsible AI, data, model risk, security, third-party risk, operations, and executive oversight. Governance must be embedded into the delivery lifecycle rather than operated as a separate policy exercise.

7.1 Governance Layers

Executive oversight

Board or executive risk committee receives AI portfolio, risk, incident, value, and remediation reporting.

AI governance steering committee

Cross-functional authority approves policies, risk tiers, exceptions, high-risk use cases, and residual risk.

Control functions

Legal, privacy, security, data, model risk, compliance, procurement, architecture, HR, and internal audit define and test controls.

Business and technical owners

Named owners are accountable for value, outcomes, data, model performance, user impact, and operational risk.

Delivery and operations

Teams use intake, impact assessment, stage gates, testing, monitoring, incident response, and retirement controls.

7.2 Enabling Technology

- AI discovery and inventory for sanctioned and shadow AI.
- Workflow-based intake, risk tiering, approvals, exceptions, and evidence collection.
- Policy and control library mapped to NIST AI RMF, ISO/IEC 42001, privacy, cybersecurity, and sector obligations.
- Model and system registry with model cards, data lineage, validation, monitoring, and change history.
- AI security controls including identity, DLP, secrets, prompt/output inspection, red teaming, and SIEM integration.
- Executive dashboard for AI value, risk, incidents, control effectiveness, third parties, and roadmap status.



8. Prioritized Recommendations

1. Establish executive governance

Approve the AI governance charter, sponsor, committee, RACI, decision rights, escalation model, risk acceptance authority, and meeting cadence.

2. Publish the policy baseline

Approve Responsible AI and Generative AI Acceptable Use policies, prohibited-use rules, exceptions, and employee obligations.

3. Build the AI inventory and intake workflow

Inventory use cases, models, agents, vendors, embedded AI, data, owners, risk tiers, lifecycle status, and approvals.

4. Implement risk tiering and impact assessment

Create risk criteria for impact, autonomy, data sensitivity, affected stakeholders, criticality, vendor dependence, and legal exposure.

5. Operationalize model and data governance

Require documentation, validation, data controls, thresholds, monitoring, change management, and retirement.

6. Implement responsible AI safeguards

Define fairness, explainability, disclosure, human oversight, appeals, harm assessment, and vulnerable-group protections.

7. Secure AI applications and agents

Apply AI threat modeling, secure SDLC, least privilege, DLP, secrets protection, prompt/output controls, monitoring, and incident response.

8. Strengthen vendor governance

Use AI-specific due diligence, contract clauses, assurance evidence, ongoing monitoring, concentration analysis, and exit plans.

9. Build audit readiness

Map obligations to controls, define evidence standards, test controls, retain records, conduct legal review, and track findings.

10. Drive adoption and continual improvement

Deliver role-based training, communications, tabletop exercises, feedback, post-launch reviews, and recurring maturity assessments.



9. 120-Day Improvement Roadmap

The workbook roadmap is appropriately sequenced from governance foundation to control validation and executive reporting. The following workstreams should be treated as an integrated transformation program with accountable owners, weekly delivery governance, dependency management, and evidence-based completion criteria.

Phase	Workstream	Improvement action	Priority	Deliverable
Days 1-30	Governance foundation	Establish AI governance committee, charter, RACI, and executive sponsor.	Critical	AI Governance Charter and RACI
Days 1-30	Inventory and intake	Create AI use case and vendor inventory, including shadow AI discovery.	Critical	AI Inventory and Intake Register
Days 1-30	Policy baseline	Publish responsible AI and generative AI acceptable use policy.	Critical	Responsible AI Policy
Days 31-60	Risk tiering	Implement AI risk tiering, impact assessment, and approval workflow.	High	Risk Tiering Methodology
Days 31-60	Data protection	Define AI data classification, privacy, retention, prompt/output controls, and DLP requirements.	High	AI Data Protection Standard
Days 31-60	Vendor governance	Deploy AI vendor risk questionnaire and contract clause checklist.	High	Vendor Risk Scorecard
Days 61-90	Model governance	Create model inventory, model card, validation, change control, and monitoring standards.	High	Model Risk Standard
Days 61-90	Fairness and oversight	Implement bias testing, explainability requirements, human oversight, and appeals process for high-risk AI.	High	AI Impact Assessment and Oversight Standard
Days 61-90	AI security	Perform AI threat modeling and integrate AI logs/events into security monitoring.	Critical	AI Threat Model and Monitoring Plan
Days 91-120	Control testing	Test AI governance controls and collect audit-ready evidence.	High	Control Test Results
Days 91-120	Incident readiness	Run tabletop for AI incident scenarios and finalize AI incident response playbooks.	High	AI Incident Response Playbook
Days 91-120	Board reporting	Launch AI governance dashboard with KPIs, KRIs, risk posture, incidents, and remediation status.	High	Board Reporting Pack

Executive Stage Gates

- Day 30: Sponsor, charter, RACI, inventory, intake, and policy baseline approved.
- Day 60: Risk tiering, data protection, privacy, and vendor governance operating for priority use cases.
- Day 90: Model governance, fairness, human oversight, AI security, and monitoring implemented for high-risk use cases.
- Day 120: Control testing, AI incident tabletop, audit evidence, residual-risk decisions, and board dashboard completed.

Recommended Executive Metrics

- Percentage of AI systems inventoried with business and technical owners.
- Percentage of AI use cases risk-tiered and approved before deployment.
- Percentage of critical/high controls with validated evidence.
- Number and age of overdue critical/high remediation actions.
- AI incidents, near misses, harmful-output events, overrides, drift alerts, and policy violations.
- Third-party AI vendors assessed and contracts updated with AI clauses.
- Training completion by role and governance workflow adoption.
- Business value delivered versus risk, cost, and control performance.





10. Conclusion

The assessment indicates a developing AI governance program with 49.7% residual risk and a Medium overall rating. The organization has meaningful strengths in data governance and selected governance practices, but the absence of approved artifacts, weak regulatory and audit evidence, and partially implemented lifecycle controls prevent a mature, defensible, and auditable operating model.

The recommended 120-day program should be executed as an enterprise transformation initiative. The first priority is to establish authority, policy, inventory, intake, and risk tiering. The second is to operationalize model, data, fairness, oversight, security, vendor, and monitoring controls. The final priority is to validate operating effectiveness, retain evidence, exercise incident response, and provide recurring executive reporting. The next assessment should replace generic responses with verified evidence, named owners, due dates, control-test results, and formally accepted residual risk.

