

PORTFOLIO PROJECT 1 OF 3

AI Governance Program Blueprint

*A NIST AI RMF, ISO/IEC 42001, and EU AI Act Implementation Plan for a
Composite Energy-Sector Enterprise*

Prepared by J Damien Scott, SRMP GRCP IAIP

801-392-9999 | jdamienscott@pm.me | jdamienscottllc.com

July 2026

Portfolio note. This document was built for a composite, fictional organization, Meridian Grid Energy Services, invented solely to demonstrate methodology. It contains no confidential or proprietary information from any real employer or client. The standards, dates, and frameworks cited are real and independently verified as of July 2026; the organization, its risk register entries, and its committee structure are illustrative.

Table of Contents

[Table of Contents](#)

[1. Purpose and Scope](#)

[2. Organizational Context](#)

[3. Governance Structure and Accountability](#)

[3.1 AI Governance Committee: Roles and Accountability \(RACI\)](#)

[4. Standards Crosswalk: NIST AI RMF, ISO/IEC 42001, and the EU AI Act](#)

[5. AI Risk Register](#)

[6. Implementation Roadmap](#)

[6.1 Phased Rollout](#)

[6.2 Cadence](#)

[6.3 AI Incident Runbook, Minimum Contents](#)

[7. Presenting This Project in an Interview](#)

[References](#)

1. Purpose and Scope

This document is a governance blueprint: a plan for how a mid-size enterprise operationalizes artificial intelligence (AI) risk management rather than a description of one. Its purpose is to show, concretely, how the four functions of the National Institute of Standards and Technology (NIST) AI Risk Management Framework, the clause structure of ISO/IEC 42001, and the obligations of the European Union AI Act translate into a governance committee, a control set, and a scored risk register that a real organization could adopt with modest adaptation.

The scope is a composite organization, Meridian Grid Energy Services (“Meridian”), a fictional regional electric and gas utility holding company with roughly 2,400 employees, invented for this project. Meridian is deploying two generative AI features: a customer-facing support assistant that answers billing questions and files outage reports, and an internal decision-support tool that summarizes grid-sensor telemetry for control-room operators. Both examples recur in Portfolio Project 2, which threat-models the customer-facing assistant in technical detail. This document addresses the governance layer above that system: who is accountable for AI risk, what controls apply, and how the organization tracks and treats residual risk over time.

A note on claim types, consistent with how I write throughout this portfolio. A fact is verifiable against a primary source, such as a standard’s published text or an official regulatory notice. Evidence supports an interpretation without settling it alone. Interpretation is my reading of what the evidence means for Meridian’s situation. Judgment is an evaluative claim about adequacy or risk. A recommendation is an action I suggest Meridian’s governance committee take. I flag each where the distinction matters.

2. Organizational Context

Meridian operates in a sector where AI governance intersects two regulatory traditions at once: utility-sector reliability and safety regulation, and the emerging AI-specific regime built from the EU AI Act, the NIST AI Risk Management Framework, and ISO/IEC 42001. A utility is also a plausible target for the kind of scrutiny that follows any customer-facing AI system: billing errors, service disruptions, and outage-related communication carry direct safety and financial consequences if an AI system gets them wrong. This context shapes the risk register in Section 5, which weights operational and safety-adjacent risks more heavily than a lower-stakes consumer application would.

Meridian’s AI footprint, for purposes of this blueprint, consists of two systems: (1) a large language model (LLM) based customer support assistant with tool access to a billing application programming interface (API) and an outage-reporting system, described fully in Portfolio Project 2; and (2) an internal grid-telemetry summarization tool used by control-room operators as a decision aid, not a decision-maker. Both systems fall within the scope of this governance program; neither is currently classified as a high-risk AI system under the EU AI Act, a classification judgment revisited in Section 4.

3. Governance Structure and Accountability

The NIST AI Risk Management Framework’s Govern function calls for clear accountability before an organization maps or measures anything. Meridian’s structure assigns a single accountable executive, a cross-functional committee, and a documented escalation path, consistent with ISO/IEC 42001’s requirement (Clause 5) for top-management commitment and defined roles within an AI management system.

3.1 AI Governance Committee: Roles and Accountability (RACI)

The table below assigns Responsible, Accountable, Consulted, and Informed status across the activities that recur most often in an AI governance program. A single accountable owner per activity avoids the diffusion of responsibility that both ISO/IEC 42001 and NIST AI RMF identify as a common governance failure.

Activity	AI Governance Committee	Chief Risk Officer	System Owner (Eng.)	Legal / Privacy	Internal Audit
Approve new AI use case for development	A	C	R	C	I
Classify AI system risk tier (EU AI Act)	C	A	C	R	I
Maintain AI risk register (Section 5)	R	A	C	C	I
Review and sign off pre-production security assessment	A	C	R	C	I
Vendor / third-party AI due diligence (Project 3)	C	A	C	R	I
Incident response for AI-related event	I	A	R	C	R
Annual management-system review (ISO/IEC 42001 Cl. 9.3)	R	A	I	I	C

R = Responsible (performs the work). A = Accountable (owns the outcome; one per row). C = Consulted (two-way input before a decision). I = Informed (notified after a decision).

4. Standards Crosswalk: NIST AI RMF, ISO/IEC 42001, and the EU AI Act

The three frameworks are complementary rather than redundant. NIST AI RMF (AI 100-1, January 2023) organizes risk management into four functions: Govern, Map, Measure, and Manage (National Institute of Standards and Technology, 2023). Its companion Generative AI Profile (AI 600-1, July 2024) extends, rather than replaces, the original framework with twelve risk categories specific to generative AI (National Institute of Standards and Technology, 2024). ISO/IEC 42001, published December 2023, is the first certifiable management-system standard for organizational AI governance, structured like other ISO management-system standards

around context, leadership, planning, support, operation, performance evaluation, and improvement (International Organization for Standardization, 2023). The EU AI Act imposes binding legal obligations, with the scope and timing determined by a system’s risk classification.

A regulatory timing update, current as of July 2026 and not reflected in older material: on May 7, 2026, European Union institutions reached political agreement on an omnibus amendment deferring the compliance deadline for most high-risk AI system obligations from August 2, 2026 to December 2, 2027. Obligations for AI embedded in regulated products, such as machinery and lifts, move to August 2, 2028. Prohibitions on unacceptable-risk systems and AI-literacy requirements took effect earlier, in February 2025, and remain in force. Meridian’s two systems described in Section 2 are not currently high-risk under the Act’s Annex III categories; this is a judgment based on the systems’ functions, not a legal determination, and Meridian’s Legal function should confirm it before any EU deployment.

NIST AI RMF Function	Core Question	ISO/IEC 42001 Clause(s)	EU AI Act Touchpoint
Govern	Who is accountable, and what policy applies before work begins?	Cl. 5 Leadership; Cl. 6 Planning	Arts. 9, 17: risk and quality management systems
Map	What is the system, its context, and its foreseeable risks?	Cl. 8.2 AI system impact assessment	Art. 9(2): risk identification and analysis
Measure	How is risk quantified, tested, and monitored?	Cl. 8.3 Data for AI systems; Cl. 9.1 Monitoring	Arts. 10, 15: data governance, accuracy, and robustness
Manage	How is residual risk treated, and how does the program improve?	Cl. 8.4 Third-party relationships; Cl. 10 Improvement	Art. 14: human oversight; Art. 26: deployer obligations

5. AI Risk Register

The register below lists ten illustrative risks drawn from Meridian’s two AI systems (Section 2) and scored on a five-point likelihood and impact scale (1 = low, 5 = high), consistent with the Measure function above. Inherent risk is likelihood multiplied by impact before controls; residual risk reflects the rating after the listed mitigation is in place. This structure, and the three-audience translation it supports (executive, engineer, auditor), follows standard risk-register practice rather than any single source.

ID	Risk	RMF Fn.	L	I	In h.	Mitigation / Control	R es .	Owner
R-01	Prompt injection causes the support assistant to disclose another customer’s billing data.	Map / Measure	4	5	20	Input fencing, output filtering keyed to account context, and per-session data-access scoping (OWASP LLM01, LLM02).	6	Eng. / Security

ID	Risk	RMF Fn.	L	I	Inh.	Mitigation / Control	Res.	Owner
R-02	Excessive agency: the assistant issues a refund or service change without human confirmation.	Manage	3	4	12	Human-in-the-loop gate on all state-changing tool calls (OWASP LLM06).	3	Eng.
R-03	Vendor-hosted foundation model changes silently, degrading accuracy on outage-reporting intents.	Measure	3	3	9	Contractual model-change notification clause; scheduled regression testing (see Project 3, Section 5).	4	Vendor Mgmt.
R-04	Telemetry summarization tool omits a safety-relevant sensor anomaly (hallucination / omission).	Measure	2	5	10	Mandatory human review before any control-room action; tool framed as decision aid, not authority.	5	Operations
R-05	Training or fine-tuning data for the assistant includes improperly retained customer PII.	Map	2	4	8	Data-minimization review at intake; DPO sign-off before any fine-tuning data set is approved.	2	Legal / Privacy
R-06	Tool-poisoning: a compromised third-party plugin manipulates the assistant's tool-call outputs.	Map / Measure	2	4	8	Allowlisted tool registry; signed tool manifests; sandboxed execution (OWASP MCP Top 10, draft).	3	Security
R-07	Unbounded consumption: a scripted user drives runaway API and token cost.	Measure	3	2	6	Per-session and per-key rate and cost caps with anomaly alerting (OWASP LLM10).	2	Engineering
R-08	System-prompt leakage exposes internal escalation logic to a curious user.	Measure	3	2	6	System-prompt segmentation; output scanning for known internal markers (OWASP LLM07).	2	Security
R-09	AI-generated code from an internal coding assistant introduces a dependency on a hallucinated package.	Manage	3	3	9	Package-existence verification gate in CI; Semgrep-class static review before merge.	3	Engineering
R-10	No documented AI incident runbook exists at time of first production incident.	Govern	3	4	12	Adopt AI-specific incident runbook (Section 6.3) prior to go-live; annual tabletop exercise.	4	Risk / Security

Column key: L = Likelihood (1–5), I = Impact (1–5), Inh. = Inherent risk score, Res. = Residual risk score after the listed mitigation.

6. Implementation Roadmap

6.1 Phased Rollout

A four-quarter rollout sequences governance formation ahead of technical control deployment, consistent with the Govern-before-Map-before-Measure ordering in the NIST framework.

Phase	Timeframe	Key Deliverables
1. Govern	Q1	Charter the AI Governance Committee; approve this blueprint; assign the RACI in Section 3.1; classify existing AI use cases against EU AI Act risk tiers.
2. Map	Q2	Complete system-level impact assessments (ISO/IEC 42001 Cl. 8.2) for both Meridian systems; populate the initial risk register (Section 5).
3. Measure	Q3	Stand up monitoring for the five detection patterns defined in Portfolio Project 2, Section 6; complete first vendor due-diligence cycle (Project 3).
4. Manage	Q4	Run the first AI incident tabletop exercise; complete the first quarterly risk-register review; brief executive leadership and the board risk committee.

6.2 Cadence

After the initial four-quarter rollout, the committee reviews the risk register quarterly, re-classifies AI system risk tiers whenever a system’s function changes materially, and repeats the incident tabletop exercise annually, mirroring the continual-improvement cycle ISO/IEC 42001 requires under Clause 10.

6.3 AI Incident Runbook, Minimum Contents

- Trigger criteria: what qualifies as an AI-related incident (data disclosure, harmful output, safety-relevant omission, cost anomaly).
- Roles: who is paged, who has authority to disable the AI feature, and who communicates externally.
- Containment step: a documented kill-switch or feature-flag procedure for each production AI system.
- Evidence preservation: prompt and response logs (hashed, not raw, consistent with Project 2, Section 6) retained for post-incident review.
- Post-incident review: root cause mapped to a specific risk-register entry and, where applicable, a MITRE ATLAS technique.

7. Presenting This Project in an Interview

This project demonstrates the ability to translate named frameworks (NIST AI RMF, ISO/IEC 42001, EU AI Act) into an operating structure a committee could adopt: a RACI, a standards crosswalk, a scored risk register, and a phased rollout. In an interview, be ready to explain why a given risk’s likelihood or impact score is what it is, and why the mitigation reduces it by that amount. That defense, not the document itself, is what a hiring panel is actually testing.

References

International Organization for Standardization. (2023). ISO/IEC 42001:2023 Information technology — Artificial intelligence — Management system.

National Institute of Standards and Technology. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0) (NIST AI 100-1).

National Institute of Standards and Technology. (2024). Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1).

OWASP Foundation. (2025). OWASP Top 10 for LLM Applications 2025. <https://genai.owasp.org/llm-top-10/>

OWASP GenAI Security Project. (n.d.). MCP Top 10 (draft). <https://owasp.org/www-project-mcp-top-10/>

European Union institutions. (2026, May 7). Political agreement on omnibus amendment deferring AI Act high-risk compliance deadlines. [Regulatory action; see official EU sources for the published text.]