

AI Governance & Risk Pack

Discovery questionnaire — score each item honestly; low scores are the point: they show where the program will hurt.

Pack	AI Governance & Risk Pack
Document	Discovery questionnaire (24 questions)
Response scales	Maturity 1–5 · Yes / Partially / No · Written
Date	2026-09

AnovaCloud — Engineering the AI-Native Enterprise. Frisco, Texas. This pack is a working instrument, not a brochure: complete it with your team, score honestly, and bring it to your discovery call. © 2026 AnovaCloud.

HOW TO USE

Run this like a workshop, not a survey

Complete it with the stakeholders named in each section's intro. One scribe, one room (or call), 90 minutes. Score from evidence — monitoring data, documents, demos — not opinions. Any item scoring 1–2 is a program risk: it needs an owner and a close date before you commit budget. Bring the scored questionnaire to your discovery call; it compresses weeks of fact-finding into one session.

1 · AI INVENTORY**AI inventory**

You cannot govern what you cannot list.

Q1. Is there a complete inventory of AI systems in use (built, bought, and shadow/employee-adopted)?

Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: the register itself, with last-review date.

Q2. Does the inventory capture data used, model/provider, owner, and deployment surface per system?

Yes

Partially

No

Evidence: sample register entries.

Q3. Are third-party and embedded-AI (SaaS features) included — not just in-house models?

Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: procurement/SaaS review coverage.

Q4. Is there a process to detect new AI adoption (intake gate, procurement hook)?

Yes

Partially

No

Evidence: intake process description.

2 · RISK TIERING**Risk tiering**

Proportional governance beats uniform bureaucracy.

Q5. Is there a defined risk-tiering model (e.g., 4 tiers by impact × autonomy)?

☐ Yes ☐ Partially ☐ No

Evidence: tier definitions.

Q6. Has every inventoried system been tiered, with high-risk systems identified?

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: tiered inventory.

Q7. Are prohibited or restricted use cases defined (and communicated)?

☐ Yes ☐ Partially ☐ No

Evidence: policy excerpt.

Q8. Is re-tiering triggered by material changes (new data, new surface, new autonomy)?

☐ Yes ☐ Partially ☐ No

Evidence: change-trigger rules.

3 · DATA & MODEL LINEAGE**Data & model lineage**

Accountability needs a paper trail.

Q9. Can you trace each AI system's training/fine-tuning data and its lineage?

☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
----------------------------	----------------------------	----------------------------	----------------------------	----------------------------

Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: lineage docs for top systems.

Q10. Are data rights (consent, license, retention) documented per dataset?

☐ Yes	☐ Partially	☐ No
------------------------------	------------------------------------	-----------------------------

Evidence: data-rights register.

Q11. Are model versions, prompts, and eval results versioned and retained?

☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
----------------------------	----------------------------	----------------------------	----------------------------	----------------------------

Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: versioning practice + retention policy.

Q12. Is there a bill of materials for third-party models/components?

☐ Yes	☐ Partially	☐ No
------------------------------	------------------------------------	-----------------------------

Evidence: AI BOM sample.

4 · TESTING & ASSURANCE**Testing & assurance**

Trust, but verify — on a schedule.

Q13. Are pre-deployment assessments required per risk tier (bias, safety, security)?

☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
----------------------------	----------------------------	----------------------------	----------------------------	----------------------------

Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: assessment templates + completion records.

Q14. Is adversarial/red-team testing performed for high-risk systems?

☐ Yes	☐ Partially	☐ No
------------------------------	------------------------------------	-----------------------------

Evidence: red-team reports.

Q15. Are eval metrics and thresholds defined per system (accuracy, fairness, grounding)?

☐ 1	☐ 2	☐ 3	☐ 4	☐ 5
----------------------------	----------------------------	----------------------------	----------------------------	----------------------------

Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: metric sheets per high-risk system.

Q16. Is post-deployment monitoring defined (drift, misuse, incident triggers)?

☐ Yes	☐ Partially	☐ No
------------------------------	------------------------------------	-----------------------------

Evidence: monitoring specs.

5 · HUMAN OVERSIGHT & INCIDENTS**Human oversight & incidents**

Someone must be able to say stop.

Q17. Is human oversight proportionate to risk tier (approval gates, override capability)?

1	2	3	4	5
---	---	---	---	---

Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: oversight matrix.

Q18. Is there an AI incident definition and response runbook?

☐	Yes	☐	Partially	☐	No
--------------------------	-----	--------------------------	-----------	--------------------------	----

Evidence: runbook + drill record.

Q19. Are affected-individual rights handled (explanation, appeal, opt-out) where required?

1	2	3	4	5
---	---	---	---	---

Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)

Evidence: rights-handling procedure.

Q20. Is there a kill-switch / rollback path for each production AI system?

☐	Yes	☐	Partially	☐	No
--------------------------	-----	--------------------------	-----------	--------------------------	----

Evidence: per-system rollback notes.

6 · REGULATORY & BOARD READINESS

Regulatory & board readiness

Governance must survive external scrutiny.

Q21. Which AI regulations apply (EU AI Act, sector rules), and is there a compliance mapping?

Evidence: regulation-to-control mapping.

Q22. Does the board (or a committee) receive AI risk reporting on a cadence?

☐	Yes	☐	Partially	☐	No
--------------------------	-----	--------------------------	-----------	--------------------------	----

Evidence: last report + cadence.

Q23. Are customer/employee disclosures about AI use in place where required?

☐ 1	☐ 2	☐ 3	☐ 4	☐ 5	Maturity 1-5 (1=Ad hoc, 2=Repeatable, 3=Defined, 4=Managed, 5=Optimized)
----------------------------	----------------------------	----------------------------	----------------------------	----------------------------	--

Evidence: disclosure inventory.

Q24. Is there an accountable AI governance owner with real escalation authority?

☐	Yes	☐	Partially	☐	No
--------------------------	-----	--------------------------	-----------	--------------------------	----

Evidence: named owner + charter.

SCORING

Read the pattern, not just the average

Average each section. Sections averaging below 3 are your critical path — sequence the project plan so those risks close first. A single 1 on a load-bearing question (downtime budget, rollback, access control) outweighs a 4 average elsewhere. Record owners and close dates below.

Lowest-scoring section:

Single biggest risk:

Owner:

Close date:
