

AI-Ready Workforce Pack

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

Pack	AI-Ready Workforce 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 · SKILLS BASELINE

Skills baseline

Start from evidence, not enthusiasm.

Q1. What % of engineers have shipped something with LLMs/APIs beyond chat (RAG, agents, evals)?

Evidence: list of shipped artifacts.

Q2. Can the team explain (correctly) embeddings, context limits, and why grounding matters?

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

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

Evidence: sampling interviews or quiz results.

Q3. Is prompt engineering / eval literacy present, or does 'AI work' mean calling an API?

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

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

Evidence: eval artifacts in repos.

Q4. Who are the 2-3 credible internal AI champions — and are they allocated, or volunteering nights?

Evidence: named people with allocated time.

2 · ROLE IMPACT MAPPING

Role impact mapping

AI changes some roles deeply and others barely.

Q5. Which roles change most in the next 18 months (support, QA, data entry, junior dev tasks)?

Evidence: role-impact map.

Q6. Are job descriptions and success metrics updated for AI-augmented work?

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

Evidence: sample updated JDs.

Q7. Is there a plan for roles that shrink (reskill paths, not just efficiency targets)?

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

Evidence: reskill path docs.

Q8. Do leaders understand what AI can/can't do well enough to set sane expectations?

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

Evidence: leadership briefing record.

3 · LEARNING PATHS

Learning paths

Training must produce artifacts, not attendance.

Q9. Are learning paths role-segmented (builders vs. leaders vs. general staff)?

☐ Yes ☐ Partially ☐ No

Evidence: path definitions.

Q10. Does enablement include hands-on sandbox builds with review — or only videos/courses?

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

Evidence: sandbox + review process.

Q11. Is there dedicated learning time (or is upskilling expected on top of delivery)?

☐ Yes ☐ Partially ☐ No

Evidence: time allocation policy.

Q12. Are internal AI patterns/playbooks captured so learning compounds?

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

Evidence: playbook repo with contributions.

4 · TOOLING & SANDBOXES

Tooling & sandboxes

People learn AI by building with it.

Q13. Is there an approved AI tool stack (models, coding assistants, agent frameworks)?

☐ Yes ☐ Partially ☐ No

Evidence: approved-tool list.

Q14. Are sandboxes available with guardrails (no production data, cost limits)?

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

Evidence: sandbox environment + policies.

Q15. Is AI-assisted development (code review, testing) normalized in the SDLC?

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

Evidence: SDLC guidance mentioning AI tooling.

Q16. Are experiments shareable (demos, internal talks) to spread learning?

☐ Yes ☐ Partially ☐ No

Evidence: demo cadence.

5 · GOVERNANCE FOR BUILDERS**Governance for builders**

Freedom within guardrails.

Q17. Do builders know the AI usage policy (data classification, approved tools, disclosure)?

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

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

Evidence: policy + attestation.

Q18. Is there a lightweight review for new AI uses (not a six-week committee)?

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

Evidence: review process + cycle time.

Q19. Are security/compliance partners embedded in enablement (not gatekeepers after)?

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

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

Evidence: joint sessions or office hours.

Q20. Are incidents and near-misses shared as learning (blameless)?

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

Evidence: incident review notes.

6 · HIRING & STAFFING**Hiring & staffing**

Build, buy, or borrow — per gap.

Q21. Which AI roles are must-hire vs. nice-to-have (ML engineer, eval specialist, AI PM)?

Evidence: prioritized hiring list.

Q22. Is the hiring bar calibrated (can interviewers assess real AI skill)?

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

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

Evidence: interview loops + rubrics.

Q23. Are contractors/partners used strategically for knowledge transfer (not just capacity)?

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

Evidence: partner SOWs with KT clauses.

Q24. Is retention risk assessed for the newly upskilled (market will notice them too)?

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

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

Evidence: retention plan.

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:
