

Agentic AI Deployment Pack

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

Pack	Agentic AI Deployment 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 · USE-CASE CANDIDACY

Use-case candidacy

Not every workflow deserves an agent.

Q1. Is the task high-volume, well-bounded, and currently painful enough to justify agent complexity?

☐ 1☐ 2☐ 3☐ 4☐ 5

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

Evidence: volume figures + current cost/pain description.

Q2. Can success be defined objectively (a checkable outcome, not a vibe)?

☐

Yes

☐

Partially

☐

No

Evidence: write the acceptance criterion in one sentence.

Q3. What is the blast radius of a wrong agent action (reversible vs. irreversible, internal vs. customer-facing)?

Evidence: list the three worst plausible failures.

Q4. Is there a human-in-the-loop step for high-stakes actions, or full autonomy from day one?

☐

Yes

☐

Partially

☐

No

Evidence: approval workflow diagram.

2 · DATA & API SURFACE

Data & API surface

Agents are only as good as the tools you give them.

Q5. Do the APIs the agent needs exist, with auth, rate limits, and idempotency documented?

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

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

Evidence: API inventory with contracts.

Q6. Is the data the agent reads fresh, governed, and access-controlled at the row/document level?

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

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

Evidence: data-access matrix.

Q7. Can the agent's actions be executed and rolled back programmatically (or are they manual-only)?

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

Evidence: per-action reversibility table.

Q8. Are tool responses structured and validated, or free-text the agent must parse?

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

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

Evidence: sample tool schemas.

3 · TOOL & INTEGRATION READINESS**Tool & integration readiness**

The plumbing decides the ceiling.

Q9. Is there a tool registry with versioning, so agent tool definitions don't drift silently?

☐

Yes

☐

Partially

☐

No

Evidence: registry or its absence.

Q10. Are long-running tool calls handled (timeouts, retries, async callbacks)?

☐

1

☐

2

☐

3

☐

4

☐

5

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

Evidence: timeout/retry policy per tool.

Q11. Can the agent environment reach production systems safely (network, secrets, egress controls)?

☐

1

☐

2

☐

3

☐

4

☐

5

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

Evidence: network/secret architecture.

Q12. Is there a sandbox/staging replica of every integrated system for agent testing?

☐

Yes

☐

Partially

☐

No

Evidence: staging inventory.

4 · EVALUATION MATURITY**Evaluation maturity**

No evals, no production. Full stop.

Q13. Is there a golden task set with expected outcomes for the target use case?

☐ Yes ☐ Partially ☐ No

Evidence: task set with N tasks and rubrics.

Q14. Are task success, tool-use correctness, and cost-per-task measured on every run?

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

Evidence: eval dashboard or its absence.

Q15. Is there regression testing on model/prompt/tool changes before rollout?

☐ Yes ☐ Partially ☐ No

Evidence: CI eval gate description.

Q16. Are failure modes catalogued (hallucinated tool calls, loops, prompt injection) with mitigations?

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

Evidence: failure-mode register.

5 · SECURITY & GUARDRAILS**Security & guardrails**

Agents are a new attack surface.

Q17. Is prompt-injection defense in place (input sanitization, instruction hierarchy, tool-output validation)?

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

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

Evidence: controls list.

Q18. Are agent permissions least-privilege, scoped per use case, and time-bounded?

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

Evidence: permission matrix.

Q19. Is PII/secrets redaction applied to prompts, tool I/O, and logs?

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

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

Evidence: redaction test results.

Q20. Has a red-team pass been run against the agent (jailbreak, data exfiltration attempts)?

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

Evidence: red-team report or scheduled date.

6 · HUMAN OVERSIGHT & OPS**Human oversight & ops**

Autonomy is earned in stages.

Q21. Are approval gates defined per action-risk tier, with named approvers and SLAs?

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

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

Evidence: approval matrix.

Q22. Is every agent run traced end-to-end (prompts, tool calls, decisions) and retained for audit?

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

Evidence: trace sample + retention policy.

Q23. Are cost guardrails set (per-task budget, kill switches on spend anomalies)?

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

Evidence: budget alerts configured.

Q24. Is there an incident runbook for agent misbehavior (halt, rollback, customer comms)?

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

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

Evidence: runbook link.

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:
