

FREE PRACTITIONER PACK

Agentic AI Deployment Pack

Five autonomy tiers, a six-question decision tree, and the oversight patterns that keep each tier honest. Autonomy is earned with measured performance — never granted by default.

Pack	Agentic AI Deployment Pack
Document	Autonomy tiers + decision tree
Use	Governance · rollout planning
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.

AnovaCloud — anovacloud.com

AUTONOMY TIERS

Tier	Definition	Human role
L0 · Advise	Agent suggests; human decides and acts.	Approve every action implicitly — the agent takes none.
L1 · Copilot	Agent drafts; human reviews, edits, and executes.	Human is the actor; agent output is a draft.
L2 · Supervised	Agent acts on pre-approved action classes; human reviews samples and exceptions.	Approval for classes, sampling review, escalation on anomaly.
L3 · Semi-autonomous	Agent acts within bounded scope; human approves only high-stakes or novel actions.	Policy-gated autonomy; human handles the long tail.
L4 · Autonomous	Agent acts end-to-end within guardrails; humans monitor aggregates and intervene on drift.	Continuous monitoring, audit, and a tested kill switch.

DECISION TREE — answer per action class, not per agent

An agent can be L1 for customer-facing messages and L3 for internal research. Score each class of action separately.

Q1 · Is the action reversible at low cost?

No → cap at L1 (copilot). Yes → continue.

Irreversible actions — sending money, deleting records, external commitments — never start above copilot.

Q2 · Could a wrong action harm a customer, employee, or the business materially?

Yes → cap at L2 with sampling review. No → continue.

Material harm includes financial loss, reputational damage, and safety consequences.

Q3 · Is the action in a regulated domain (credit, hiring, medical, legal, safety)?

Yes → cap at L2 and involve compliance before any expansion. No → continue.

Regulators care about the decision process, not the tool — document the human's role explicitly.

Q4 · Can the agent's work be verified automatically (tests, rules, a second model)?

Yes → L3 is reachable with policy gates. No → stay at L2 or below.

Verification is what converts oversight from a bottleneck into a control.

Q5 · Is there a named human with capacity to review exceptions within the SLA?

No → do not exceed L2. Yes → continue.

Oversight without capacity is theater. Staff the checkpoint before you open the tier.

Q6 · Do you have 4+ weeks of measured performance at the current tier?

No → hold the tier. Yes → consider one tier up, with rollback criteria written first.

Tier promotion is a change request: new tier, new guardrails, new monitoring — and a way back.

OVERSIGHT PATTERNS

Approve-before-execute

The agent proposes; nothing runs until a human approves. Use for irreversible or high-blast-radius actions. Keep the approval queue visible and SLA-bound.

Review-after-execute (sampling)

The agent acts; humans review a sampled fraction plus all anomalies. Sampling rate starts high and drops only as measured quality holds.

Escalation triggers

Hard rules that force human review: confidence below threshold, novel situation, policy-flagged content, cost above limit, or customer-impacting outcome.

Stop & override

Any reviewer can pause the agent and take over mid-task. The override path must be tested like a fire drill — not discovered during an incident.

Record the chosen tier per action class in the Agent Permission Matrix, and review tiers quarterly — or after any incident, whichever comes first.