

AI Agent Autonomy & Guardrails Checklist

A production checklist of the guardrails every autonomous AI agent needs — action budgets, approval gates, sandboxing, validators, kill switches — plus an autonomy maturity model.

By Syed Husnain Haider Bukhari — AI engineer & founder, Revolutionary Technologies · www.husnainbukhari.com

AI Agent Autonomy & Guardrails Checklist

This is the working checklist I run every production AI agent through before it takes a single autonomous action on real traffic. It expands on the ideas in my article on constrained autonomy into a deliverable you can actually execute against: risk tiers, a per-action policy template, a nine-layer controls stack, a five-stage maturity model, and a pre-launch checklist. Use it as a gate, not as reading material — an agent should not ship until you can tick every applicable box.

1. Classify Every Action by Blast Radius

Before you talk about autonomy at all, inventory every action your agent can take and assign each one a risk tier. Autonomy is granted per action, never per agent, and the tier is what determines the default level. Reversibility and reach are the two axes that matter most.

Tier	Definition	Examples	Default autonomy
Green	Reversible, internal, cheap	Read CRM, enrich a record, summarize a call, draft text	Autonomous with post-hoc review
Amber	Reversible but customer-visible or costly	Update a deal stage, schedule within known availability, tag a contact	Autonomous within tight rails
Red	Irreversible, external, or financial	Send email to a real person, charge or refund, delete a record, post publicly	Human approval required

If you cannot confidently place an action in a tier, treat it as Red until proven otherwise. The cost of over-constraining a Green action is a slightly slower workflow; the cost of under-constraining a Red one is an incident.

2. The Autonomy Spectrum

Every action sits at one of five autonomy levels. Write the level next to each action in your inventory.

- **Level 0 — Suggest only.** Agent proposes, human executes. No autonomous action.
- **Level 1 — Execute with pre-approval.** Agent prepares the exact payload; a human approves before it runs.
- **Level 2 — Execute within tight rails.** Agent acts inside a narrow allowlist of tools, arguments, and value ranges; anything outside escalates.
- **Level 3 — Execute with post-hoc review.** Agent acts freely within a category; a human reviews a sample afterward.
- **Level 4 — Self-directed within a budget.** Agent chooses actions and stopping point, bounded only by hard budgets and kill switches.

The question is never how autonomous can this agent be. It is what is the worst thing this agent can do before a human notices, and can I live with that.

3. The Nine-Layer Controls Stack

Each layer catches a different class of failure. Skipping a layer chooses which incident you will have. Build them roughly in this order.

1. **Action budgets and step limits.** Max steps, tool calls, wall-clock time, USD per run, and consecutive errors before abort. The cheapest control and the biggest money-saver.
2. **Tool allowlists and argument constraints.** Not just which tools exist, but which tool with which arguments against which resources. Scope tools to the task at hand.
3. **Human-in-the-loop approval gates.** For every Red action, pause and present a legible diff — exact action, payload, and stated reason — so a human can decide in seconds.
4. **Sandboxing and least privilege.** Run code and web browsing in isolation, with scoped credentials, egress rules, and no standing access. Process untrusted content away from action-taking tools.
5. **Output schemas and validators.** No free text flows into an action. Every tool call passes a strict schema plus business-rule validators (refund $\leq$ charge, no past-dated bookings, discount $\leq$ ceiling).
6. **Rate limits and cost caps.** Per-agent and per-tenant, enforced outside the model. A runaway loop should trip a cap and page a human, not rack up a bill.

7. **Observability and tracing.** Full trace per run: inputs, each reasoning step, every tool call with args and results, tokens, latency, outcome. Replayable.
8. **Kill switches and circuit breakers.** A no-deploy global halt plus self-tripping breakers on error rate, reject rate, and cost per run.
9. **Evals as the change gate.** Representative tasks with known-good outcomes and adversarial traps. Ship prompt, model, and tool changes only when the numbers hold.

4. Per-Action Autonomy Policy Template

Encode the policy declaratively so it is reviewable by a non-engineer. Adapt this shape to your agent.

```
{
  "agent": "your-agent-name",
  "budgets": {
    "max_steps": 15,
    "max_tool_calls": 30,
    "max_wall_clock_seconds": 120,
    "max_usd_per_run": 0.40,
    "max_consecutive_errors": 3
  },
  "actions": {
    "datasource.read": { "autonomy": "level_3", "approval": false },
    "record.update": { "autonomy": "level_2",
      "allow": { "field_in": ["stage","tag","note"] } },
    "email.send": { "autonomy": "level_1", "approval": true,
      "validators": ["recipient_known", "template_approved"] },
    "payment.refund": { "autonomy": "level_0" }
  },
  "circuit_breakers": {
    "reject_rate_pct": 20,
    "cost_per_run_usd": 1.00,
    "on_trip": "halt_and_page"
  },
  "kill_switch": "flags.agent_enabled"
}
```

When a stakeholder asks to loosen a control, the change is a single reviewable line — exactly what you want to take into a risk conversation.

5. The Autonomy Maturity Model

You advance a stage because you have evidence, not because time has passed. Define the entry condition for each stage before you start.

Stage	What the agent does	Entry evidence required
1 — Shadow	Proposes only, takes no action	Proposal-vs-human agreement above your bar on a real sample
2 — Assisted	Executes only with per-action approval	High approval rate, low edit rate, understood rejection reasons
3 — Supervised	Autonomous on Green/Amber, approval on Red	Stable daily traces, sensitive circuit breakers not tripping
4 — Bounded	Autonomous across most scope under hard budgets	Clean post-hoc sample review over a sustained window
5 — Delegated	Self-directs within a budget	Narrow, low-blast-radius domain only; rarely justified

Refuse to skip stages. A great demo is Stage 1 evidence at best. The model gives you shared language to say "not yet" without sounding like you are stalling.

6. Defenses Most Teams Skip Until an Incident

Prompt injection

The moment your agent reads untrusted content — a web page, an inbound email, a scraped profile, a support ticket — that content can carry instructions. Isolate untrusted input from your action-taking tools, keep a strict tool allowlist, and never let retrieved text silently become a system instruction. Add adversarial injection cases to your eval suite so a regression is caught before it ships.

Irreversible actions

Gate every irreversible action behind a dry-run and an approval. Distinguish test records from live ones structurally, not by convention. Prefer designs where the agent stages a batch that a human releases, rather than firing one-by-one with no chance to intervene.

Compounding errors

Give the loop a natural stopping point. Step budgets and consecutive-error limits turn a confused agent into a fast, cheap failure that escalates, instead of a slow, expensive thrash that drifts further from the goal.

7. Pre-Launch Guardrails Checklist

- Every action is inventoried, tiered, and assigned an autonomy level.

- Hard budgets exist for steps, tool calls, time, cost, and consecutive errors.
- Tools are allowlisted per task with argument validation.
- Every Red action has a legible human approval gate.
- Untrusted content is processed in isolation from action tools.
- All tool calls pass strict schemas plus business-rule validators.
- Per-tenant rate limits and cost caps are enforced outside the model.
- Full, replayable traces are emitted for every run.
- A no-deploy global kill switch and self-tripping circuit breakers are live.
- An eval suite with adversarial cases gates every change.
- The starting maturity stage is set to Shadow or Assisted, not Bounded.

8. Instrument These Metrics Before You Grant More Autonomy

You cannot responsibly turn the autonomy dial on a feeling. Each notch of freedom should be justified by numbers you were already tracking. These are the metrics I watch, and the direction each needs to move before I advance an agent to the next maturity stage.

Metric	What it tells you	What good looks like
Approval rate	How often humans accept the agent's proposed action unedited	High and stable before moving an action off approval
Edit rate	How often humans change the payload before approving	Low; a high edit rate means the agent is not ready
Breaker trip rate	How often circuit breakers halt the agent	Rare and explainable, not routine
Cost per successful run	Whether budgets are tuned or the agent thrashes	Predictable, well inside the cap
Escalation rate	How often the agent hits a rail and hands off	Trending down as rails are tuned to real traffic
Eval pass rate	Whether changes regress behavior or adversarial defenses	Green on the full suite, including injection cases

Set the target for each metric in advance, in writing, as the entry condition for the next stage. Deciding the bar after you have the numbers is how teams rationalize shipping something that was not ready.

9. A Two-Week Rollout Sequence

Here is the sequence I use to take an agent from nothing to supervised autonomy without skipping evidence. Treat the durations as minimums, not deadlines.

1. **Days 1–2.** Inventory actions, assign risk tiers, and write the per-action policy. No code yet — this is the design gate.
2. **Days 3–5.** Build the envelope first: budgets, allowlist, validators, tracing, kill switch. Wire the agent to run in Shadow mode with actions stubbed.
3. **Days 6–8.** Run Shadow on real inputs. Compare proposals to human decisions. Build the eval suite from the traces you collect, including the failures.
4. **Days 9–11.** Move to Assisted. Turn on approval gates for every action. Measure approval and edit rates on live traffic.
5. **Days 12–14.** If the metrics clear the bar, move Green and Amber actions to supervised autonomy with breakers sensitive and traces reviewed daily. Red actions stay gated.

If any stage fails its bar, you stay there and fix the agent — you do not push forward on schedule pressure. The sequence is deliberately unglamorous, and that is the point: durable agents come from the boring envelope, not the clever prompt.

10. Common Failure Modes to Avoid

- Adding controls after the first incident instead of before the first deploy.
- Confusing a happy-path demo with production readiness.
- Granting broad tool access for convenience and discovering the blast radius during an incident.
- Treating observability as a nice-to-have, so every failure is a mystery.
- Tuning autonomy on vibes because there is no eval suite to measure against.

If you want a second set of eyes on your agent's controls before it goes live — or help taking an agent from an impressive demo to a system you can trust with customers, money, and data — [work with Syed Husnain Haider Bukhari](#). Teams shipping outbound and prospecting agents can see many of these guardrails in practice inside [ProLeads](#). Ship the constrained version first; you can always turn the dial up, but you cannot un-send the email.