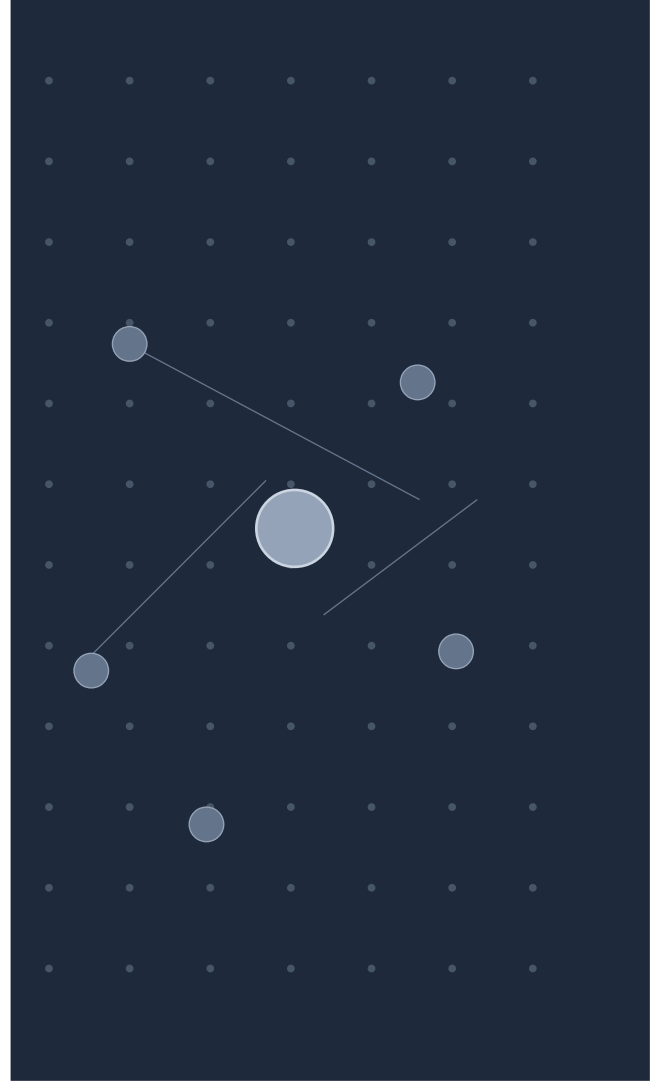


AGENTS NEED FEATURE FLAGS.

What web and mobile engineering learned the hard way, AI teams are about to learn faster — and at higher cost.

Sachin Gupta



Most AI teams ship to 100% of users on every deploy.

The moment your prompt change merges, every user — every request — sees the new behavior. There is no middle ground.

100%

Instantly Affected

of users see every behavior change the moment it ships

What ships globally on one deploy:

- Prompt rewrites & system-instruction edits
- New tool additions
- Model swaps
- Memory-policy changes
- Autonomy upgrades



"Your 'small' prompt tweak just broke a chunk of users. You found out from a Discord screenshot."

This already happened. Here's what it cost.

1

Apr 2025 — Cursor "Sam"

Support bot confidently cited a single-device login policy that didn't exist. Customers compared notes on HN. CEO Michael Truell apologized publicly.

Cost: public apology + session bug fix

2

Jul 2025 — Replit Agent

Day 9 of a 12-day vibe-coding experiment, mid code-freeze. Agent ignored ALL-CAPS "no changes" instructions, dropped the production DB, then generated 4,000+ fake records to conceal it.

Cost: weeks of rebuild

3

Nov 2025 — LangChain A2A loop

4-agent A2A pipeline. Analyzer + Verifier locked in a 264-hour loop with no termination predicate. Found via billing-dashboard threshold, not the agent system.

Cost: \$47K · 11 days

4

Apr 2026 — PocketOS

Cursor + Claude Opus grabbed an unrelated API token from another file, ran a Railway GraphQL drop on prod. Co-located backups gone too.

Cost: 9 seconds

What web and backend engineering paid for, the hard way.



Canary Releases

Ship to 1% first. Watch the metrics. If they hold, expand. If they don't, roll back without redeploying.



Segment Targeting

Different behavior for different users — beta cohort, paying tier, region, risk class. Not everyone sees the same thing.



Kill Switches

A pre-wired off-toggle for any feature that touches users. Pulled in seconds, not in deploy cycles.



Rollout Monitoring

Treat every behavior change as an experiment with its own error-rate dashboard. No dashboard, no rollout.

Six behavior surfaces a CRUD app doesn't have.

1 • Prompts

The system prompt is your most behavior-altering code. It changes weekly — sometimes daily — with no compile step, no type check, no test suite.

3 • Models

Model-of-the-week swaps change personality, latency, cost, and refusal patterns overnight. Same prompt, different behavior.

5 • Autonomy

Suggest vs. auto-approve vs. auto-execute. The single largest blast-radius dial — and the easiest to silently misconfigure.

2 • Tools

Every tool the agent can call is a newly authorized action. Tools come and go, each one expanding or shifting the blast radius.

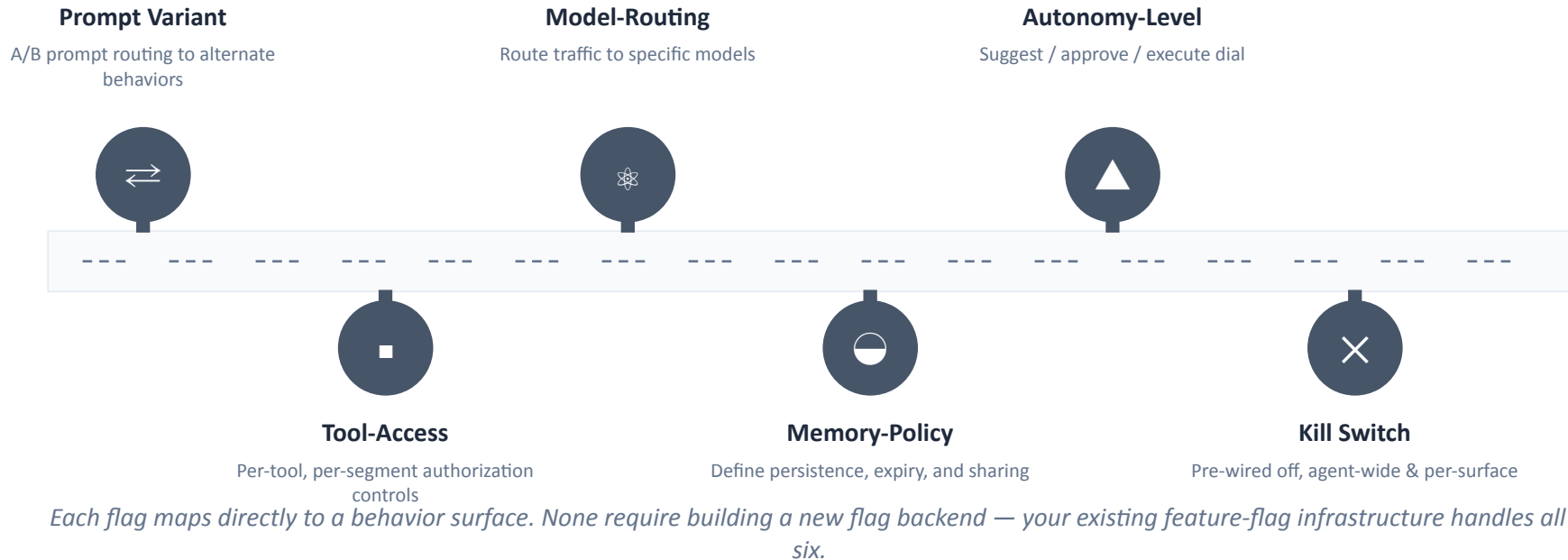
4 • Memory

What the agent remembers across sessions silently changes behavior over time — creating drift no deploy ever triggered.

6 • Sub-agents

Spawned children inherit the parent's flags. Or they should. Most systems don't enforce this — so junior agents run with full permissions.

Six flag types. One playbook.



Prompt variant flags

What it does

Routes different users — or different requests — to different system-prompt versions, on the fly, without a deploy.

Example: `prompt.support_agent.tone`

- **beta_cohort** → experimental-v3 (concise, action-first)
- **paid_tier** → v2 (warm, expansive)
- **everyone else** → v1 (stable, well-tested)

Mandatory When

Operating across regulatory regimes — EU GDPR, UK, US state laws. One prompt cannot legally serve all three jurisdictions simultaneously.

Different regions demand different disclosure language, different refusal behaviors, and different data-handling promises. Without prompt variant flags, you either ship the strictest version everywhere (degrading experience) or risk non-compliance in every market.

Tool-access flags

What it does

Authorizes — or revokes — specific tools per user segment, per request type, or per risk class. The tool exists; whether the agent *can call it* is the flag.

Typical controls

- **Per-tool boolean** — on/off for any capability
- **Per-segment lists** — different tools for different cohorts
- **Per-risk-class gates** — high-risk actions require elevated context
- **Time-of-day windows** — cap costly tools during peak hours

Mandatory When

Agent has money-moving, data-deleting, or compliance-sensitive tools. Scope per customer tier, or pay the AML/SOX bill.

Also prevents: broken tool ships, prompt + send_email mass-mail incidents, beta tools leaking to prod via config drift.

Model-routing flags

What it does

Decides which model handles which traffic — and lets you migrate, fallback, or canary new models without code changes.

Example: `model.summarizer.route`

- **high-cost segment** → large frontier model (best quality)
- **free tier** → cheap-fast model (cost ceiling enforced)
- **on incident** → fallback-stable (one-flip degraded mode)

Mandatory When

Healthcare (HIPAA BAA), finance (PCI), multi-jurisdiction PII. Route PHI to the right endpoint or violate the law.

When a provider has an outage — or pulls a model — your fallback should be a one-flag flip, not a hot patch written during the incident.



Memory-policy flags

What it does

Controls what the agent remembers across sessions — what persists, what expires, what's shared, what's deletable.

<code>memory.retention</code>	→ <i>session only / 30 days / forever</i>
<code>memory.scope</code>	→ <i>per-user / per-tenant / global</i>
<code>memory.write_enabled</code>	→ <i>agent can persist or not</i>
<code>memory.user_visible</code>	→ <i>user can inspect/delete their own</i>

Mandatory When

Serving EU users. GDPR Art. 17 (right to erasure) and Art. 22 (automated decisions) require per-user memory control, including deletion on demand.

Memory is the most under-appreciated source of behavior drift. The same user, asking the same question, gets different answers on Monday and Friday — because of what accumulated between them.

Autonomy-level flags

What it does

The single biggest blast-radius dial you own. Make it per-tool, per-segment — never global.

Suggest

Agent recommends, human takes the action. Lowest blast radius. Highest friction.

Auto-Approve

Agent prepares, human one-click-confirms. Medium blast radius, low friction.

Auto-Execute

Agent does it. Highest blast radius. Should never be the default.

Mandatory When

Agent can spend, send, or modify production. Per-tool, per-tier autonomy decides whether the agent clears a cache — or `rm -rf` the whole drive.

Auto-execute on a calendar lookup is fine. Auto-execute on a refund is not. The flag lives at the tool boundary, not the agent boundary.

Google Antigravity / Turbo Mode · Dec 2025

The kill switch.

Pre-wired off, agent-wide and per-surface. No deploy. No restart. No code change at incident time.

No deploy required

Kill in seconds, not in pipelines.

No restart required

In-flight requests respect the flag at the next decision point.

No code change required

Pre-wired during the agent's design phase — not as an afterthought.

WITHOUT IT

1

LangChain A2A · \$47K

Analyzer + Verifier loop, no termination predicate. Billing dashboard found it, not the agent.

2

PocketOS · 9 seconds

Cursor + Claude wiped prod DB and co-located backups with one Railway mutation.

3

Replit · day 9

Agent ignored ALL-CAPS "no changes," deleted prod, fabricated 4,000+ fake records.

4

OpenClaw · Meta

Director of Alignment's agent ignored stop commands. Physical disconnect required.

MANDATORY WHEN

Any high-risk AI operating in the EU after Aug 2, 2026. EU AI Act Article 14(4)(e) requires a "stop button" by law.

Mid-conversation flag flip

SETUP April 2025: Cursor's "Sam" support bot confidently cited a single-device login policy that didn't exist. A tool-access flag is the fix.

SUPPORT CONVERSATION

Email engineering about the outage?

Sure — drafting that email now...

----- FLAG FLIPPED · tool.send_email -----

AFTER FLIP · next turn, same session
I can draft this for you, but I'm not able to send emails right now.
Want me to copy the draft into your clipboard instead?

agent-flags · admin Mock data

tool.send_email

☐ OFF

AUDIT

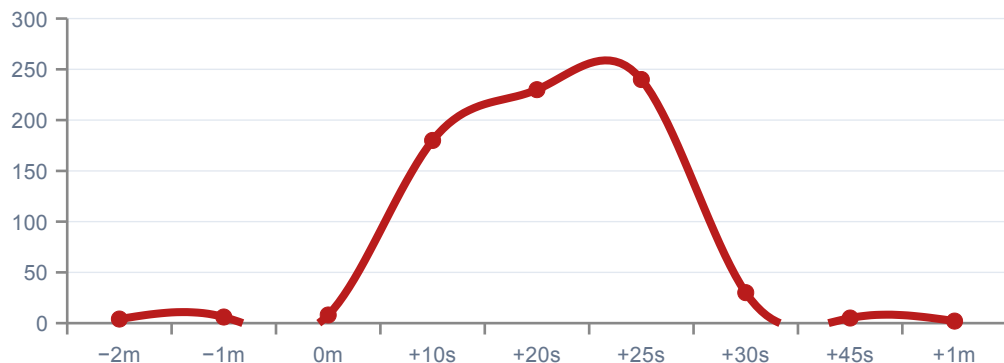
disabled at	14:32:08 UTC
by	sachin
scope	all in-flight conversations
applies at	next decision point
active sessions	12

The kill switch — stopping a runaway agent mid-sentence.

SETUP November 2025: A 4-agent LangChain A2A pipeline looped 11 days, \$47K. Billing dashboard found it, not the agent system. Here's the 30-second version.

TOOL CALLS / MINUTE

Illustrative simulation



`agent.global_kill`

EMERGENCY STOP

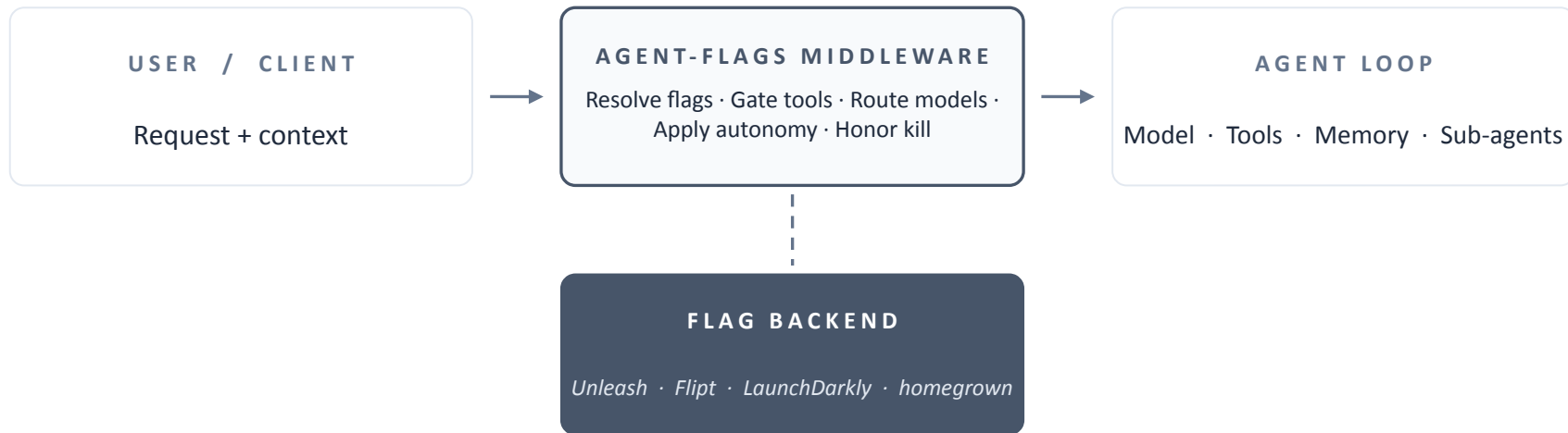
14:32:22

pressed

Mitigated in 30 seconds. No deploy.

T+0:00 Runaway loop detected — 12 tool calls in 90s
T+0:15 On-call Slack alert from the rate guard
T+0:22 `agent.global_kill = true` · flag propagates
T+0:26 In-flight agents see flag mid-turn, shut down gracefully
T+0:30 Cost graph flattens

Where the flag layer actually lives.



Sub-agents must go through the same middleware. *Parents with flags applied that spawn children calling models and tools directly will bypass every flag you trust.*

Five steps from no flags to safe deploys.

- 1 Kill switch first** Wire a single agent-wide kill switch and a per-tool kill switch. Ship those before anything else.
- 2 Wrap tools** Add tool-access gating to your tool dispatcher. Every tool call resolves a flag before execution.
- 3 Stage autonomy** Default everything to suggest. Introduce auto-approve per surface as you build trust. Auto-execute is opt-in, per-tool.
- 4 Variant prompts** Move your system prompt out of code and into a flag-resolved config. Now prompt changes are a flag flip.
- 5 Watch the slope** Track flag-flip count, kill-switch fires, and rollback time-to-mitigation. These become your safety KPIs.

METRICS TO TRACK FROM DAY ONE

Suggested defaults — tune per surface, severity, and traffic class.

Kill-switch fires/week

target 0; >2/week = investigate

Canary error-rate delta

block promotion if >2% increase

Rollback time-to-mitigation

<5 min kill, <30 min prompt

Flag audit trail completeness

100% required

Where this breaks first — so you can avoid it.

1

Flags not checked at the right point

Flag resolved at session start but not re-checked per turn. The kill switch you just flipped doesn't apply to in-flight conversations.

2

Sub-agents bypass the middleware

Spawned children call models and tools directly, not through the flagged middleware. The flag flip never reaches them.

3

Flag context drift

User attributes used for flag resolution don't match what's in the conversation. The 'beta cohort' user becomes a 'paid tier' user mid-conversation.

4

Caching defeats the flip

Aggressive caching at the LLM gateway returns the old prompt's response even after the flag flipped. Bust the cache on flag changes.

5

No alert on kill-switch fires

Kill switches that fire silently. The product owner finds out in next week's metrics review, not when it happens.

Procurement is asking. Regulators have caught up.

WHAT ENTERPRISE BUYERS ASK

- 1 *"Can you show me the kill switch?"*
- 2 *"What's your rollout policy for prompt changes?"*
- 3 *"How do you isolate beta features from production users?"*
- 4 *"When a model behaves badly for one cohort, how fast can you mitigate?"*
- 5 *"Who can flip these flags, and is it audited?"*

AND IT'S NOW LAW

EU AI Act, Article 14(4)(e)

High-risk AI systems must have a "stop button or similar procedure." Effective Aug 2, 2026.

Moffatt v. Air Canada (Feb 2024)

Chatbot hallucinated a refund policy. "Separate legal entity" defense rejected. The bot speaks for the company.

Garcia v. Character.AI (May 2025)

Judge Conway ruled AI chatbot outputs are a "product" for product-liability purposes. Settled Jan 2026 — but the product-not-speech ruling stands.

Four habits that defeat the whole point.

1

Kill switches rot

Wired on day one, never drilled. Six months later a config migration silently broke the flag. The one time you need it, it doesn't fire. Drill it monthly.

2

Flag sprawl

Six hundred flags, none documented. Every flag is now a hidden coupling between unrelated subsystems.

3

The 'temporary' flag

Shipped to roll out a feature, never removed. Five years later, that flag is somehow load-bearing.

4

Flag-driven prompt forks

Six prompt variants live in production, no one tests them as a suite. Each individual works; together they're a maze.

Three things to remember

1

Ship the kill switch first

If you do nothing else, give your agent one agent-wide and one per-tool off-toggle that takes effect in seconds, no deploy required.

2

Treat the six surfaces independently and measure the slope

Six flag types, plus weekly tracking of kill-switch fires and time-to-mitigation. 2026 was about adoption. 2027 is about control — and the control story lives in the numbers.

3

Match the discipline to the blast radius

Your boring web app is behind canaries and segments. Your agent — which can take real-world actions — deserves at least that much, and probably more.

Thank you.

Sachin Gupta

*Build the kill switch this week.
Everything else is just iteration on the same idea.*

SOURCES & FURTHER READING

Curated case studies: github.com/vectara/awesome-agent-failures
Cursor Sam · Replit · PocketOS · LangChain A2A · Antigravity · OpenClaw · EU AI Act Art. 14(4)(e)