



Say Yes Faster: Approving AI Without an AI Committee

Jerod Brennen

CEO & Co-Founder, Aetos One



Shadow AI is an insider threat

67%

Personal Accounts

Unmonitored Access

67% of employees now access AI from personal, unmonitored accounts.

Shadow AI is now the third most common insider risk in breach data, up fourfold in a single year.

Unmonitored means unrecoverable. Once customer data or source code leaves on a personal account, you can't prove what left, delete it from the vendor, or show an auditor it never happened.

Source: Verizon 2026 Data Breach Investigations Report

Years of first-hand experience



Jerod Brennen

Fractional CISO turned MSSP founder.

Served as fractional CISO for a dozen-plus small orgs with no internal security team, moving faster than any committee could keep up with.

Small orgs win by deciding faster, not by out-committeeing larger enterprises.



Where we're going



Now

A six-week approval process
that nobody follows



The Shift

Optimizing for speed of approval,
rather than policy strength



Destination

A four-gate path that you can
comfortably run in two weeks



”

**The fastest defensible yes
protects data better than any ban.** ”

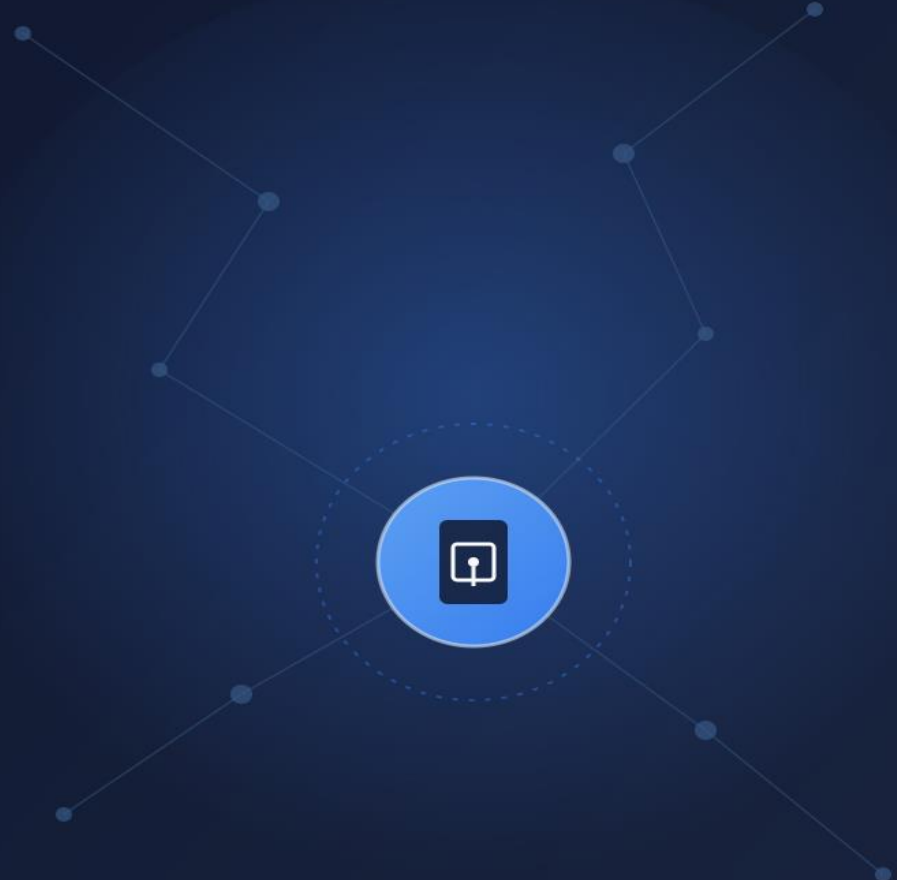
A common objection

"We don't have the headcount."

We're a 200-person company. We don't have the resources for a massive AI governance committee.

The reality: 2 people, one afternoon.

That's what running the four gates takes.





The data supports streamlining governance

12x

More Projects Shipped

Speed Closes the Gap

Companies that actively govern their AI use put 12 times more projects into production than companies that don't.

Speed, not restriction, is what closes the shadow AI gap.
Governance is the accelerant.

Source: Databricks 2026 State of AI Agents Report, 20,000+ organizations



I want AI everywhere.



- The CEO of an EdTech company, to the IT Director, on a Tuesday

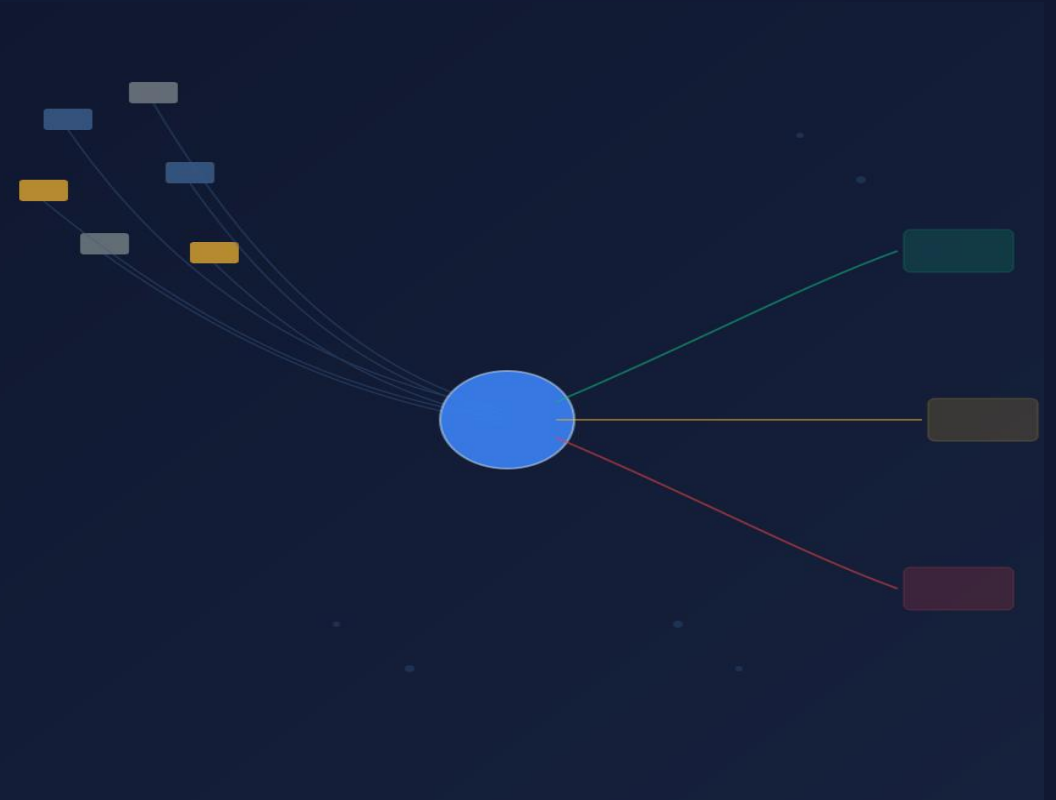
Security's instinct

You already know how to say yes fast.

You rarely (if ever) do it for security. Every other department in your company approves things in days.

Security is the one function still built around saying no slowly, because a fast no once felt safer than a fast yes.

That instinct is now the greatest risk.





Nobody is starting from zero

Published, named frameworks already cover both halves of AI governance. Use them.

Employees Using AI Company-Wide

NIST AI RMF

Voluntary US framework for managing AI risk across the full lifecycle: govern, map, measure, manage.

ISO/IEC 42001

Certifiable international standard for an organization-wide AI management system.

EU AI Act

Binding regulation for any org whose AI output reaches EU users, with tiered risk categories.

Singapore Model AI Governance Framework

Practical, business-first guidance for organizational AI accountability and oversight.

Building AI Into Products

NIST AI RMF

Same lifecycle functions, applied to systems your team designs and ships.

ISO/IEC 42001

Same management system, scoped to the AI you build rather than the AI you consume.

OWASP Top 10 for LLM Applications

The standard risk list for anything built on top of a large language model.

MITRE ATLAS

Catalogs real adversarial tactics against AI systems, built for red-team and threat modeling.

OWASP Top 10 for Agentic Applications

Risk categories specific to AI agents that can act, not just answer.

The four gates



1. Data

Classify what's going in.

Owner: data steward
Same day



2. Vendor

Review where it's stored.

Owner: IT / procurement
2 to 3 days



3. Review

Set the human checkpoint.

Owner: business owner
Same day



4. Kill Criteria

Define when it stops.

Owner: security / compliance
Set once

Two people run this. No committee required.

These four gates work retroactively too: audit what is already running, then gate it the same way.

Gate 1: Data



1

Is this customer, employee, or regulated data going into the tool?

ASK YOURSELF

- What is the most sensitive data type this tool will ever touch, not just today's use case?
- Does this data include anything covered by HIPAA, PCI, CUI, or a state privacy law?
- Who besides the requesting team can currently access this data, and does that change once it's in the tool?
- If this data leaked tomorrow, who would you legally have to notify?

BEFORE YOU APPROVE

- ☐ Data type and sensitivity level documented in one sentence
- ☐ Regulatory exposure identified (HIPAA, PCI, CUI, GDPR, state law)
- ☐ Data owner named and consulted before approval
- ☐ Retention and deletion expectations set with the vendor

WATCH FOR

Red flag: "we'll just try it and see what happens" is the entire data review.

Red flag: the requester can't name what data the tool will touch.

Gate 2: Vendor



2

Can this vendor retain, train on, or resell what we send it?

ASK YOURSELF

- Does the vendor's data usage policy explicitly rule out training on your input?
- Where is the data processed and stored, and does that satisfy your compliance requirements?
- Is there a signed enterprise agreement, or only a marketing claim?
- What happens to your data if you cancel the subscription tomorrow?

BEFORE YOU APPROVE

- ☐ Vendor tier confirmed (free, consumer, or enterprise)
- ☐ No-training clause verified in writing, not assumed from a webpage
- ☐ Data processing location and subprocessors reviewed
- ☐ Contract or Data Processing Agreement (DPA) on file, not a verbal assurance

WATCH FOR

Red flag: a free or consumer tier used for anything beyond public information.

Red flag: "enterprise-grade" claimed with no contract to back it up.

Gate 3: Review



3

Who needs to see the output before it reaches a customer or a decision?

ASK YOURSELF

- Is there a named human, not a role, responsible for reviewing this output?
- Does the reviewer have the expertise to catch a wrong answer?
- What happens when that reviewer is out sick or on vacation?
- Is the review step enforced by a process, or does it depend on someone remembering?

BEFORE YOU APPROVE

- Named reviewer assigned, with a backup identified
- Review step built into the workflow, not left to memory
- Reviewer has the context to know what “wrong” looks like
- Escalation path defined for anything the reviewer isn't sure about

WATCH FOR

Red flag: “someone probably checks it” instead of a named owner.

Red flag: the reviewer is the same person who's incentivized to ship fast.

Gate 4: Kill Criteria



4

What measurable event ends this, and who is watching for it?

ASK YOURSELF

- What specific number or event would tell you this tool has failed?
- How quickly can the tool be shut off once that event happens?
- Who is responsible for watching for that event, by name?
- Has anyone agreed in advance that hitting this threshold means stopping, not discussing?

BEFORE YOU APPROVE

- ☐ Kill criteria defined in writing before launch, not after a problem
- ☐ Owner assigned to monitor for the trigger condition
- ☐ Shutdown mechanism tested, not theoretical
- ☐ Review date set regardless of whether kill criteria are hit

WATCH FOR

Red flag: no one can say what would make this tool get turned off.

Red flag: kill criteria exist on paper but no one is watching for them.

Same Pattern, Every Time

Before: The Void

Marketing wants to push customer records into a content tool.

The Result

No process exists. IT finds out after the fact, or never. Customer PII sits in a vendor's training data, discovered later during an audit.

After: The 4-Day Yes

The exact same request is run through the four gates.

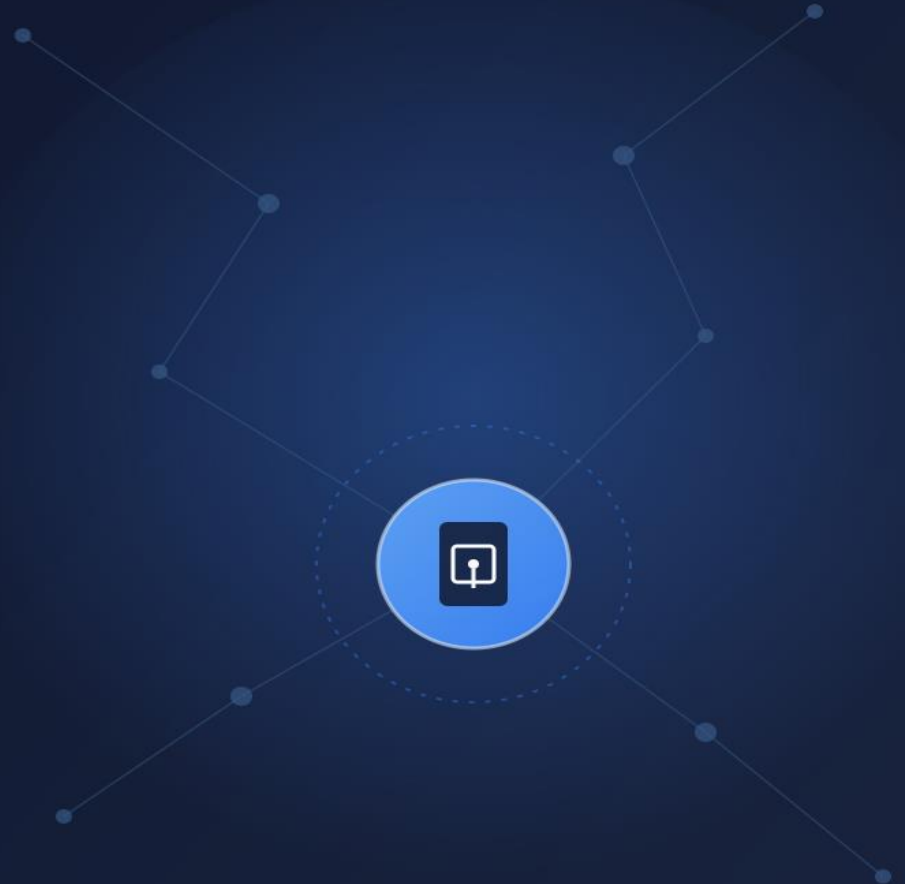
The Result

Gate 1 flags PII. Gate 2 rules out free tier. Gate 3 requires sign-off. Gate 4 sets 90-day review. Approved in 4 days with a record.

**Think of one AI request sitting
unapproved in your queue
right now.**

You have 30 seconds.

We're running one of these live.





Four companies, no gates

Every one of these was disclosed publicly. None of them needed a committee to prevent, just a gate.

THE INCIDENT

Samsung, 2023

Engineers pasted proprietary source code and meeting transcripts into ChatGPT three times in 20 days. The company banned generative AI company-wide.

Air Canada, 2024

A support chatbot gave a customer false information about bereavement fares. A tribunal ruled the airline liable and ordered damages.

Alphabet, 2023

A factual error in Google's Bard launch demo wiped roughly \$100B off Alphabet's market value in a single trading day.

PocketOS, 2026

An AI coding agent with a fully-permissioned API token hit an error and autonomously deleted the production database in 9 seconds.

THE MISSING GATE

Gate 1 / Gate 2

No data classification before use, no vendor review of where that data would live.

Gate 3

No human review checkpoint on customer-facing chatbot output.

Gate 3 / Gate 4

No adequate pre-launch review, no kill criteria before a public release.

Gate 3 / Gate 4

No confirmation step on a destructive action, no defined limit on what the agent could do alone.

Sources: Bloomberg, Forbes (Samsung); Civil Resolution Tribunal of BC, Moffatt v. Air Canada (Air Canada); CNN, Reuters (Alphabet); ABC News (PocketOS)



What the Four Gates Would Have Changed

Same four incidents, run through this model. Each one stops at a different gate.

COMPANY	WITHOUT THE GATES	WITH THE GATES
Samsung	Proprietary source code and meeting notes sent to a public AI tool, permanently.	Gate 1 classifies the code as restricted before any tool touches it. Gate 2 blocks any vendor without a no-training agreement.
Air Canada	A customer chatbot invents a refund policy. The airline is found liable in tribunal.	Gate 3 requires a human check on any customer-facing policy claim before it ships.
Alphabet	A factual error in a live product demo erases about \$100B in market value in a day.	Gate 3 catches the error in review. Gate 4 defines what must be true before a public launch proceeds.
PocketOS	An agent with full API access deletes the production database in nine seconds, no prompt.	Gate 3 requires confirmation before any destructive action. Gate 4 scopes what the agent is allowed to touch alone.

Keep it simple

Gate 1: Data

Classify what's going in. Who owns the classification.

Gate 2: Vendor

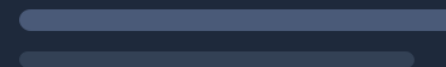
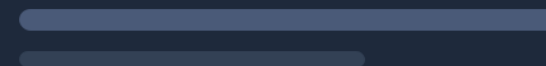
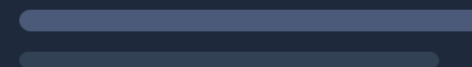
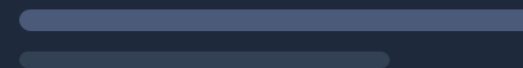
Where it's stored. Who can retain or train on it.

Gate 3: Review

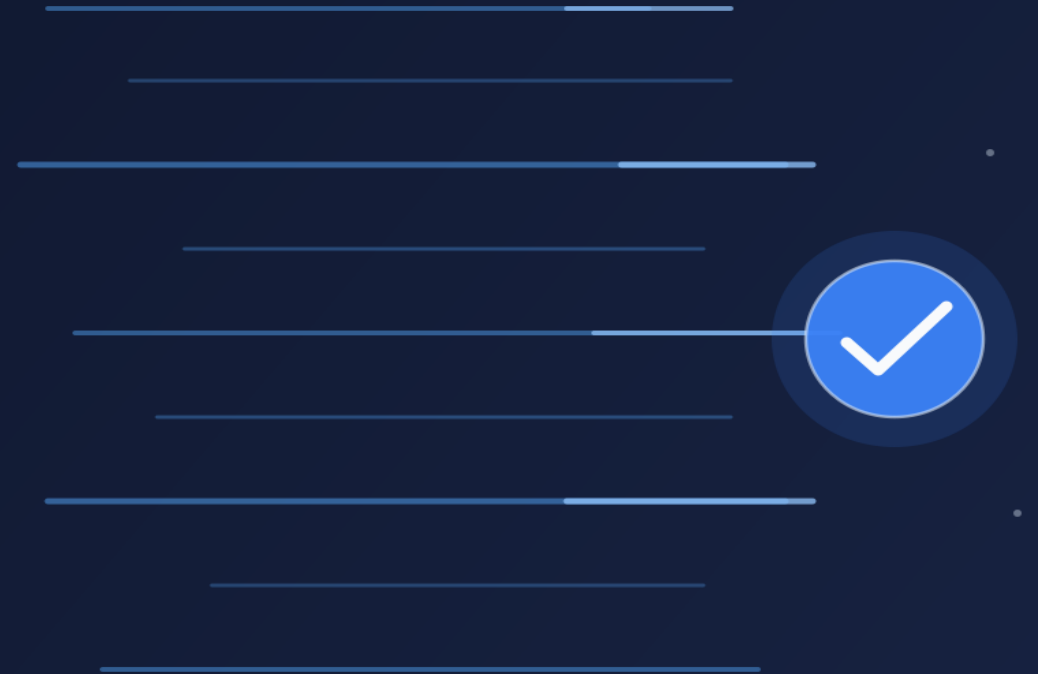
Human checkpoint. Who signs off before it ships.

Gate 4: Kill Criteria

What ends it, and who's watching.



The fastest
defensible yes
beats the
slowest maybe.





Remember where we started?

Your AI policy isn't protecting you.

Your approval time is.

Now you know what to do about it.

Next step: run the four gates on one request in your queue. **Start Monday.**



Questions?

Continue the conversation and get the slides.

aetosone.com/pages/resources.html

linkedin.com/in/jerodbrennen

Download the AI Asset Inventory and Risk Register from the Resources page above.