

Healthcare AI governance lives in the workflow, not the policy binder

Ask most healthcare organizations whether they have AI governance, and the answers are all over the map.

A [HFMA and Eliciting Insights survey](#) of health system executives found 71% have already deployed AI solutions in clinical, finance, or revenue cycle functions, but **just 18% have a mature governance structure** and a fully formed AI strategy to match.

Among the organizations that do report having governance, the picture gets more complicated. [RSM's 2026 middle-market AI survey](#) found 94% of healthcare respondents have governance controls in place, but **13% only apply those controls after something has already gone wrong**, and 9% apply them inconsistently depending on the team or use case. One healthcare IT executive in RSM's focus groups summed up the gap:

“We’re trying to be compliant without knowing what compliance is.”

A written policy can say the right things and still not tell you what happens the moment an AI system runs into something it shouldn’t handle alone. That’s the actual test: not what the document promises, but what the system does when it counts.



The moment that actually matters in healthcare AI

Picture a routine outreach call: an AI agent working through a health plan’s gap-closure list, checking in on members who are due for a screening or a follow-up. Partway through the conversation, the member mentions their blood pressure reading, and it’s elevated.

- **Ungoverned:** the reading gets logged. It shows up, accurately, in a weekly report. Nothing routes it anywhere in the moment.
- **Governed:** the exact same reading triggers an escalation the instant it’s captured. A clinician sees it, with the full encounter and member claims history attached, not a summary written after the fact. If this member’s history shows a prior cardiac event, that context is already in view, not discovered later.

That clinician sees it in time to act, not react, which can be the difference between catching a problem early and discovering it in an ER visit three weeks later. AI didn’t get smarter between these two versions. The workflow around it did.

That’s the impact governance makes. It’s also the simplest test for whether a healthcare organization’s AI governance is real in practice: not what the policy says should happen, but what the system actually does the moment something needs human clinical intervention.

One AI principle, three very different situations

In care navigation, the same design, with deterministic clinical logic plus a human escalation path, shows up differently depending on what the AI is actually doing:

- **Symptom triage.** A member describes their condition, and AI guides the conversation from there. The next step comes from an established clinical protocol, not AI’s own judgment, and a clinician signs off before anything gets recommended.
- **Gap closure.** A member mentions something clinically relevant, like the blood pressure reading above, and it escalates immediately instead of waiting for the next report cycle.
- **Post-discharge engagement.** After leaving the hospital, a member reports their symptoms are worse than expected. Instead of just getting logged alongside everyone else’s updates, it goes straight to a clinician.

Three different moments in a member’s care, one consistent rule across them all: AI can run the conversation, but it doesn’t decide what happens when a situation requires clinical judgment.

Engage with AI. Decide with clinicians.

That rule breaks down into three distinct roles, and keeping them distinct is the mechanism behind everything above:

- **Engage.** Conversational AI handles the warm, member-friendly outreach and intake, the part designed for members to actually finish the conversation.
- **Assess.** Deterministic protocols, not the model’s own reasoning, drive the clinical logic. This step decides whether something needs human clinical intervention.
- **Decide.** A licensed clinician reviews and makes the decision, with the full picture in front of them rather than a summary reconstructed after the fact.

AI never renders the clinical decision itself. It runs the conversation, hands off the moment assessment calls for it, and lets an expert decide.

Why AI governance can’t be retrofitted in healthcare

It’s tempting to think governance can be layered onto an existing AI system after the fact, a review process added here, an audit trail bolted on there. In practice, that’s much harder than building the escalation path in from the start. Retrofitting requires knowing, in advance, every place a silent gap might exist in a system that was never designed to surface them.

That highlights the argument for governance as engineering rather than governance as policy. A policy can say “escalate clinically significant findings.” Only the underlying architecture can guarantee that an elevated reading, mentioned once, in passing, in the middle of an unrelated conversation, actually reaches someone right away, instead of waiting for the weekly report where it would otherwise surface.

What “built in” actually looks like

A workflow that’s actually governed, rather than governed on paper, includes engineering receipts:

- **Full interaction logging.** Every interaction is logged and tied to a durable encounter ID, end to end, not summarized after the fact.
- **Per-turn model tracking.** The AI model and version behind every conversational turn is recorded, so any response can be traced back to its source.
- **Human review trail.** Clinician review, approval, and override actions are captured with the encounter itself, proving the human was actually in the loop rather than asserting it.
- **Append-only audit store.** The record lives inside the compliance boundary and can only be added to, never rewritten, so it can’t be cleaned up after the fact.

These are the markers of a verifiable governance structure, the kind a trustworthy clinical AI partner can demonstrate.

The Pager HealthSM approach to healthcare AI governance

Pager Health’s [care navigation](#) and [wellness](#) platforms apply this standard everywhere. [Five guardrails](#) apply to every interaction, engineered into the platform itself rather than added after a missed escalation:

- Data privacy and security
- Clinical accuracy and hallucination prevention
- Ethical and bias mitigation
- Human in the loop
- Operational and lifecycle governance

With these guardrails in place, informed by more than a decade of healthcare AI experience, Pager Health connects more than 26 million individuals across the US and Latin America to the right care at the right time, helping healthcare organizations close care gaps, reduce total cost of care, and improve health outcomes. And when members reach out directly, that access holds up in practice: the median time to connect with a nurse is 6 seconds.¹

The takeaway

Whether a healthcare organization has no governance structure at all, or a policy that looks solid on paper, the actual test is the same: what happens the moment an AI system encounters something it shouldn’t handle alone. That’s a question only the workflow can answer.

See how Pager Health’s guardrails work in practice →

¹ Pager Health Nurse Advice Line data, August 2026 year to date.

Solutions

[Pager HealthSM Navigator](#)
[ReallyWellSM – Member Login](#)
[Enterprise 360](#)

Customers

[For Payers](#)
[For Providers](#)
[For Employers](#)

Resources

[Blog](#)
[White Papers](#)
[Case Studies](#)

About

[Company](#)
[Careers](#)
[Press](#)

Contact

[Contact Us](#)