
ISG-AMERICAS HEALTHCARE · LEADERSHIP SERIES · NO. 3

Who Signs?

Accountability, Governance, and the Leadership Gap in Clinical AI

American health systems have deployed artificial intelligence into clinical workflows faster than they have decided who is answerable when it is wrong. That is not a technology gap. It is a leadership vacancy, and most organizations have not yet noticed they have one.

Lee Fossey

US Sector Head, Healthcare · Executive Managing Partner · ISG-Americas

From the boardroom to the bedside · August 2026

CONTENTS

What this paper covers

I. Executive Summary

II. The Deployment Is Already Done

III. The Accountability Vacuum

IV. What Goes Wrong, and Why It Is Not Rare

V. The Regulatory Floor Is Rising

VI. The Chief AI Officer Question

VII. Why Clinicians Override

VIII. The Executive Profile That Barely Exists

IX. The Board Agenda

X. What Good Looks Like

XI. Conclusion: Who Signs

XII. References

A note on sources. Every figure in this paper is attributed. Where studies disagree, both readings are given. Nothing here is intended as legal advice, and organizations should take counsel on their own exposure.

Executive Summary

Three quarters of United States health systems now run at least one artificial intelligence application, up from three fifths a year earlier.¹ The Food and Drug Administration has authorized more than fourteen hundred AI-enabled medical devices, and authorized more of them in 2025 alone than in the agency's entire first decade of doing so.² Four in five physicians report using AI professionally, a figure that has more than doubled in three years.³ Whatever anyone believes about the wisdom of the pace, the argument about whether healthcare will use this technology has been settled by the people already using it.

What has not been settled, in a striking number of organizations, is the far less exciting question of who is answerable when it goes wrong.

This paper argues four things. First, that the accountability gap is real, structural, and largely unexamined, because American malpractice doctrine currently routes the consequences of an algorithmic error to the clinician at the bedside rather than to the executives who procured, validated and mandated the tool.⁴ Second, that the failure modes are documented rather than hypothetical, and the two most instructive cases involved tools that were already deployed at scale before anyone checked.^{5,6} Third, that the regulatory floor is rising quickly enough that governance built reactively will be built under pressure and at cost.^{7,8,9} And fourth, that the scarce resource in all of this is not data scientists. It is executives who combine clinical credibility, technical literacy, and the organizational standing to halt a deployment their own board has already celebrated.

**Every health system in America can tell you who bought the algorithm.
Very few can tell you who is answerable when it is wrong.**

The recommendation that follows from all of this is unglamorous. Name one clinical executive who owns clinical AI outcomes, give that person genuine authority to stop a deployment without seeking permission, govern models the way the organization already governs any other clinical intervention, and hire for the judgment to refuse rather than the fluency to explain. Organizations that do this will capture the benefit of the technology, which is real and is already measurable in clinician burnout data.¹⁰ Organizations that fill the org chart and hope will discover what they bought at the worst possible moment.

75%

of US health systems used at least one AI application in 2026, up from 59% in 2025¹

1,451

AI-enabled medical devices authorized by the FDA as of March 2026²

47%

of physicians rank increased oversight as the single regulatory

PART II

The Deployment Is Already Done

It is worth establishing the scale precisely, because a great deal of governance discussion still proceeds as though this were a question about the future.

Adoption at the organizational level

Survey work published in 2026 puts adoption of at least one AI application at seventy five percent of United States health systems, against fifty nine percent the year before.¹ Half of systems are now running three or more applications, and the number of organizations implementing or planning three or more rose by two thirds year over year.¹ Broader hospital surveys place overall use closer to eighty percent. Clinical note taking and ambient listening lead the field at sixty eight percent adoption. Roughly seven in ten non-federal acute care hospitals have integrated predictive AI into the electronic health record, and around three quarters of hospitals use AI-assisted diagnostic tooling somewhere in radiology.¹

Adoption at the device level

The regulatory record tells the same story from a different direction. The FDA's public database listed roughly nine hundred and fifty AI-enabled devices in August 2024, more than twelve hundred and fifty by July 2025, and one thousand four hundred and fifty one by March 2026.² The agency authorized three hundred and thirty one such devices in 2025, the most in its history, against six in 2015.² Around three quarters of the authorized estate sits in radiology.

Milestone	Figure
FDA AI-enabled devices authorized, August 2024	~950
FDA AI-enabled devices authorized, July 2025	~1,250
FDA AI-enabled devices authorized, March 2026	1,451
Authorized during calendar 2025 alone	331
Authorized during calendar 2015	6
Share of authorized devices in radiology	~76%

The category that regulation has not reached

Here is the detail that matters most for governance, and it is routinely missed. As of March 2026, no FDA-authorized device used generative AI or was powered by a large language model.² Yet generative tooling is precisely what has spread fastest through clinical settings, in the form of ambient documentation, drafted patient messages, summarization and appeal letters. A large and growing share of the AI that clinicians actually touch every day therefore sits outside the device authorization pathway entirely. An executive who assumes FDA clearance is doing the safety work is, for much of the deployed estate, assuming something that is not true.

The practical consequence

Two quite different regimes are operating at once. Imaging and predictive devices largely run through a regulatory pathway with defined evidence requirements. Generative and workflow tooling largely does not. Most organizations govern them, if at all, through the same committee, at the same cadence, with the same assumptions. That is the seam where problems will surface first.

The benefit is real, and should be said plainly

None of this is an argument against deployment. The clearest evidence sits in the ambient documentation literature. A 2025 study following ambulatory clinicians across six health systems recorded burnout falling from fifty two percent to thirty nine percent after thirty days of ambient scribe use.¹⁰ Mass General Brigham reported a twenty one point absolute reduction in burnout prevalence, from fifty three percent to thirty one percent.¹⁰ A randomized trial at UW Health found around thirty minutes less documentation time per provider per day, while a larger 2026 study across five academic medical centers found more modest gains of roughly sixteen minutes saved per eight hours of patient care.¹⁰ The honest reading is that the well-being effect is consistent and substantial, and the time-savings effect is real but smaller and more variable than the marketing suggests. Notably, the largest benefits accrued to clinicians who were coached on how to work with the tool rather than left to discover it alone, which is itself a leadership finding rather than a technology one.

PART III

The Accountability Vacuum

Ask, in a leadership meeting, who is answerable when a deployed model contributes to a patient being missed. The scene that follows is familiar to anyone who has sat through it. The vendor points at the contract and the intended-use statement. Legal points at the clinician, because the clinician made the call. The clinician points out they were instructed to use the tool and were not shown its performance characteristics. Information technology points out, correctly, that they implemented what they were

asked to implement. Everyone in the room is being entirely sincere, and there is no one at the end of the chain.

What the law currently does with that ambiguity

American malpractice doctrine resolves the ambiguity in a way most executives have not internalized. Liability continues to rest on the standard of the reasonable physician under similar circumstances, and courts judge the physician's conduct whether or not AI was involved.⁴ Analysis of electronic health record litigation suggests courts focus on how the clinician interpreted and acted on the software's output, applying the principle that the physician-patient duty of care is non-delegable.⁴ The legal system is not, in the main, asking whether the algorithm failed. It is asking what the clinician did.

The exposure runs in both directions at once, which is the part clinicians find hardest to accept and the part they are most right about. A physician may be found liable for relying on an erroneous AI recommendation, and may equally be exposed for failing to use an available and accurate AI tool.⁴ Meanwhile health systems and physicians are generally exposed under negligence theories but unlikely to face products liability, while algorithm developers are more plausibly exposed to products liability but unlikely to be reached in negligence.⁴ There is no doctrine assigning shared responsibility to an AI system, even where its recommendation directly shaped the decision.⁴

The organization procures the tool, validates it or does not, mandates its use, and sets the workflow. The clinician carries the consequence. That is the arrangement, whether or not anyone chose it.

Why this is a governance failure rather than a legal curiosity

Read the previous paragraphs as an executive rather than as a lawyer and the implication is uncomfortable. An institution can make every material decision about a clinical AI tool, which vendor, which population, which threshold, whether to validate locally, whether to monitor for drift, whether the alert is dismissible, and still find that the risk lands on the individual who was handed the output at two in the morning. That is not a stable arrangement. It is also, quietly, an ethical problem, because the party best positioned to prevent the harm is not the party carrying it.

The physician workforce has noticed. In the American Medical Association's 2026 survey, eighty seven percent of physicians identified assurances that they would not be held liable for errors made by AI models as important to their trust in these tools, and eighty six percent named medical liability coverage.³ Eighty five percent want to be consulted on or responsible for AI adoption in their practice.³ This is not technophobia. It is a workforce reading the liability structure accurately and asking to be included in decisions whose consequences they will personally carry.

The question that resolves it

There is a straightforward diagnostic. For every model running in clinical workflow, can the organization name the individual, not the committee, who is accountable for its performance in this population, and does that individual have the authority to switch it off? Where the answer is yes, most of the governance follows naturally, because a named person with real authority will insist on validation, monitoring and an off switch out of self-interest. Where the answer is no, no amount of policy documentation compensates, because policy without a name attached is a description of what everyone assumes someone else is doing.

PART IV

What Goes Wrong, and Why It Is Not Rare

Two cases deserve to be understood in detail by every healthcare executive, not because they are unusual, but because both involved tools already deployed at scale before anyone independent checked them.

CASE ONE

The sepsis model that was widely deployed and poorly validated

In 2021, researchers at Michigan Medicine published an external validation of a widely implemented proprietary sepsis prediction model in *JAMA Internal Medicine*, covering 27,697 patients across 38,455 hospitalizations.⁵ The model returned a sensitivity of thirty three percent, a specificity of eighty three percent, a positive predictive value of twelve percent, and an area under the curve of 0.63.⁵ In plain terms, it missed roughly two thirds of sepsis cases, and of the patients it did flag, close to nine in ten did not have the condition.

The accompanying editorial was titled around the importance of external validation, and that is the lesson.⁵ The model was not obscure. It was embedded in a widely used electronic health record and running in hospitals across the country. Its ability to identify sepsis had simply not been adequately evaluated on the populations where it was operating. Nobody behaved maliciously. The tool was procured, switched on, and trusted, and the checking step was assumed to have happened somewhere upstream.

CASE TWO

The population health algorithm that encoded unequal access

In 2019, Obermeyer and colleagues published an analysis in *Science* of a risk prediction algorithm applied to millions of patients.⁶ At any given risk score, Black patients were considerably sicker than White patients. Correcting the disparity would have raised the share of Black patients identified for additional support from 17.7 percent to 46.5 percent.⁶

The mechanism is the instructive part. The algorithm was not fed race. It predicted future health care costs as a proxy for future health needs, and because unequal access means less is spent caring for Black patients at equivalent levels of illness, the proxy carried the inequity forward at scale. By conventional accuracy metrics the model performed well. It was accurately predicting the wrong thing. As the authors put it, the choice of a convenient proxy for ground truth is itself a major source of algorithmic bias.⁶

What the two cases have in common

Neither failure was a coding error, and neither would have been caught by asking a vendor whether their model was accurate. In the first case the model was validated somewhere other than where it ran. In the second the model was accurate against its stated target and the stated target was wrong. Both are questions of judgment about clinical purpose and local population, and both sit precisely where clinical leadership and technical literacy have to meet. Neither is answerable by a procurement process, an IT function, or a data science team working alone.

The distinction that governance must encode

A model that performed well elsewhere, on somebody else's patients, has not been validated for your organization. It has been advertised to it. Validation is a local activity, repeated over time, because the population drifts, the workflow changes, and the model does not know that either has happened.

PART V

The Regulatory Floor Is Rising

Organizations that treat AI governance as optional are making a bet on a landscape that has been moving quickly for three years.

Federal transparency requirements

The Office of the National Coordinator's HTI-1 final rule introduced the first requirements of their kind for transparency in predictive algorithms inside certified health IT.⁷ Certified modules must support thirty one source attributes for predictive decision support interventions, against thirteen for conventional evidence-based ones.⁷ Those attributes cover what data trained the intervention, how it should be used and maintained, how it performs on validity and fairness metrics both in testing and against local data, and they extend to social determinants and health equity data sources.⁷ The organizing standard asks whether an intervention is fair, appropriate, valid, effective and safe.⁷

Read as an executive rather than a compliance officer, that rule does something quite specific. It makes the information necessary to govern these tools available, which removes the most common excuse for not governing them. Once the performance characteristics are disclosable, choosing not to look becomes a decision rather than a limitation.

Accreditation moves in

In June 2025 the Joint Commission and the Coalition for Health AI announced a partnership, and in September 2025 released guidance on the Responsible Use of AI in Healthcare.⁸ The guidance is non-binding, and covers governance, privacy and transparency, data security, safety event reporting, and risk and bias assessment.⁸ It recommends formal AI usage policies and a governance structure drawing on compliance, information technology, clinical programs, operations and data privacy.⁸ Governance playbooks were flagged to follow, and in May 2026 a voluntary certification in the responsible use of AI in healthcare was released.⁹

The direction of travel here is unambiguous to anyone who has watched accreditation bodies operate. Voluntary guidance becomes voluntary certification, voluntary certification becomes a competitive expectation, and competitive expectation becomes a survey question. Organizations that build governance now will be describing something they already do. Organizations that wait will be building it against a deadline, which is both more expensive and less likely to produce something that works.

The states are not waiting

During 2025, forty seven states introduced more than two hundred and fifty AI-related bills, and thirty three were enacted.¹¹ For any system operating across state lines this produces an obligation that is neither uniform nor stable, and it is not a problem a single compliance officer can hold in their head. It is a reason to build one internal standard that meets the strictest applicable requirement, rather than fifty local answers.

Voluntary guidance becomes certification. Certification becomes expectation. Expectation becomes a survey question. The only variable is whether you build it deliberately or under a deadline.

The Chief AI Officer Question

The instinctive organizational response has been to create a role. Estimates suggest that between thirty and fifty percent of healthcare organizations will have hired or be hiring an AI officer by the end of 2026, against five to ten percent previously.¹² Searches for healthcare AI leadership roles ran at more than double the prior year's pace through the first four months of 2026.¹² Kaiser Permanente named a Chief AI Officer in the first quarter of 2026, and Providence formalized a role it had operated functionally since 2024.¹² New titles have proliferated, including Director of AI Governance and Responsible AI Lead, and the authority attached to them varies enormously.¹²

Creating the role is, in the right structure, a good decision. But the title is not the variable that determines whether it works. Three structural questions decide that, and they are worth asking before the search begins rather than after the hire disappoints.

One. Where does the role report?

A role reporting through technology will be measured on delivery, integration and adoption, because that is what technology functions are properly measured on. A role reporting through clinical leadership will be measured on patient outcomes and safety. Both are legitimate structures and they produce different behavior under pressure. The question to ask honestly is what happens when a deployment the organization has publicly committed to shows a safety signal three months in. Under the first structure the instinct is to fix and continue. Under the second it is to pause and evaluate. Neither instinct is wrong in every case, but only one of them is what the institution will wish it had when the case is serious.

Two. Does the role carry clinical credibility?

Medical staff assess new authority quickly and largely on whether the person has carried clinical risk. An executive who has never been personally responsible for a patient outcome can be respected, resourced and entirely correct, and will still struggle to persuade a skeptical department to change practice. This is not a fairness question, it is an operational one. The organization needs the clinical body to actually adopt what is deployed, and adoption follows credibility.

Three. Can the role stop something?

This is the decisive question and the one most often left unanswered. Can this person halt a live deployment without first obtaining consensus from the executives who sponsored it? If the answer is no, the organization has created an advisory function and should be honest with itself that it has done so. Advisory functions produce excellent memoranda and cannot prevent anything.

There is a failure mode specific to creating a dedicated AI function, and it is subtle enough to catch capable organizations. Once AI becomes a named portfolio belonging to someone, the chief medical officer and chief nursing officer can reasonably conclude it is no longer theirs. Ownership contracts to one office, the operators who live with the consequences disengage from decisions about it, and the institution ends up with a specialist who is accountable in name and an executive team that has quietly stepped back. The strongest structures treat clinical AI as a clinical intervention owned by clinical leadership, with specialist capability supporting it, rather than as a technology portfolio held to one side.

PART VII

Why Clinicians Override

A persistent executive complaint is that clinicians are slow to adopt. The data does not support it. Eighty one percent of physicians now report using AI professionally, more than double the thirty eight percent recorded in 2023, and more than three quarters believe it improves their ability to care for patients, up from sixty five percent.³ This is not a resistant workforce. It is an adopting workforce with specific and well-founded conditions.

The precedent that shapes every new alert

Clinical staff are not evaluating new tools in a vacuum. They are evaluating them after a decade of decision support that cried wolf. Meta-analysis of drug interaction alerting puts physician override prevalence at around ninety percent, with reported ranges across systems and contexts running from roughly half to almost all alerts.¹³ One analysis found that of more than thirty eight thousand alerts classified as very severe, eighty eight percent were overridden.¹³

That behavior is usually described as alert fatigue, framed as a human failing to be corrected with training. It is more accurately described as a rational response to a system that has repeatedly demanded attention without earning it. Clinicians learned, correctly, that most alerts did not warrant the interruption. Any new tool inherits that learned skepticism on arrival, and no amount of enthusiasm at the executive level transfers to the point of care without something more substantial behind it.

Clinicians are not asking whether the tool works. They are asking whether anyone will stand behind them if they follow it and it is wrong.

What they are actually asking for

The AMA data is specific about the conditions. Forty seven percent of physicians rank increased oversight as the single regulatory action that would most increase their trust, making it the leading

answer.³ Eighty seven percent point to data privacy assurances and the same proportion to not being held liable for errors made by AI models.³ Ninety two percent want more education and training.³ Separately, sixty one percent reported alarm at payers' expanding use of AI in prior authorization, which is worth noting because it demonstrates that clinician concern is not about the technology in general but about who is wielding it and to what end.³

Every item on that list is a leadership deliverable rather than a technical one. Oversight, liability clarity, education, and consultation before adoption are all things an executive team either provides or does not. Trust in clinical AI is not a change management problem to be solved with a communications plan. It is an output of visible accountability, and it has to come from someone the clinical body already believes.

PART VIII

The Executive Profile That Barely Exists

Everything above converges on a single scarce person. It is worth describing that person precisely, because organizations that search for the wrong profile will hire quickly and then discover the gap remains.

The four attributes

Attribute	What it means in practice
Clinical credibility	Has personally carried clinical risk and is believed by the medical staff. Can walk into a departmental meeting and be heard as a peer rather than as administration.
Technical literacy	Can interrogate a model's performance characteristics rather than accept a vendor's summary. Understands the difference between discrimination and calibration, knows why a strong area under the curve can coexist with a useless positive predictive value, and asks which population the validation was performed on.
Operational standing	Sits close enough to the executive table that a decision to pause is a decision, not a recommendation. Has budget or workflow authority that makes the pause enforceable.
Willingness to refuse	Will stop a deployment that leadership has already announced, and will accept the professional cost of doing so. This is a character attribute and it is the rarest of the four.

Why this profile is scarce

Nothing in the standard healthcare executive career produces it. A physician executive career develops clinical credibility and operational standing while rarely developing quantitative depth. An informatics

or analytics career develops technical literacy while rarely conferring the clinical standing that persuades a skeptical medical staff. Technology leadership develops delivery capability and almost never confers authority to override a clinical decision. The combination has historically had no career track, no credential, and no obvious feeder role. It is being assembled by individuals, one at a time, largely by accident.

The willingness to refuse is scarcer still, and it deserves saying plainly. Stopping a project the board is proud of carries a real cost to the person who does it, and organizations rarely reward it visibly. Most executives, being sensible, will find a way to raise a concern that does not require them to be the one who halted the initiative. This is why the authority has to be structural rather than personal. If stopping a deployment depends on individual courage, it will happen too late.

Where these people are actually found

- **Physician executives who took quantitative work seriously.** Chief medical officers, chief medical information officers and physician leaders who have run informatics or quality analytics alongside clinical responsibility. This is the most common origin and the most reliable.
- **Clinical quality and patient safety leaders.** Underestimated as a source. They already run event reporting, root cause analysis and intervention withdrawal, which is precisely the machinery clinical AI governance requires. The technical layer is the addition, and it is the easier half to add.
- **Nursing executives with informatics depth.** Frequently overlooked, and often the strongest on adoption because nursing workflow is where most clinical AI actually lands.
- **Payer-side and risk-bearing medical directors.** Have spent careers making population-level decisions on imperfect data with real consequences attached, which is an unusually good preparation for the judgment this role requires.

What connects all four is not credentials. It is that each involves having made consequential decisions under uncertainty and having been personally answerable for them. That is the thing being hired for, and it is not visible on a résumé keyword search.

The hiring error to avoid

The most common mistake is hiring for fluency instead of judgment. The candidate who explains the architecture most impressively is usually not the candidate who will decide correctly whether to trust a model with a patient at three in the morning. Fluency is comparatively abundant and is becoming more so. Judgment under clinical uncertainty is not, and it does not interview as well because it tends to sound cautious.

The Board Agenda

At some point a board will ask what is running in the organization and how anyone knows it is safe. A significant number of institutions could not currently answer. The following questions are ordered so that the first three will reveal, on their own, whether the rest are worth asking.

1. **How many models are currently running in clinical workflow, and who maintains that list?** A surprising number of organizations cannot produce the inventory. If the list does not exist, nothing downstream of it does either.
2. **For each one, name the individual accountable for its clinical performance.** Not the committee. A person, with a title.
3. **Does that person have authority to switch it off without seeking consensus?** If not, the accountability is nominal.
4. **Which models were validated on our own population, and when?** Vendor validation on another institution's patients is not validation here.
5. **How is performance monitored now that our population has changed since deployment?** Drift is not an edge case, it is the default condition.
6. **What is the process for withdrawing a model, and has it ever been used?** A withdrawal process that has never been exercised is a document, not a capability.
7. **Which deployed tools sit outside FDA authorization?** For most systems this will include the generative tooling clinicians use most.²
8. **How do clinicians report an AI-related safety concern, and where does that report go?** If it routes to a vendor ticket queue, it is not safety reporting.
9. **What proportion of alerts from each tool are overridden, and what have we done about the outliers?** Override rate is the most honest available measure of whether clinicians find a tool worth their attention.¹³
10. **What are we contractually owed by vendors on performance disclosure and post-deployment monitoring?** Most legacy contracts say considerably less than executives assume.
11. **If a patient were harmed tomorrow with a model involved, who from this organization would answer for it?** If the honest answer is the clinician at the bedside, the board now knows what it needs to fix.

A board that can get satisfying answers to all eleven is governing this technology. A board that stalls at the second question has learned something more valuable than any presentation would have told it.

What Good Looks Like

Organizations handling this well tend to do the same six things, and none of them require exceptional resources.

1. One name, not one committee

A single clinical executive owns clinical AI outcomes. Committees advise them. Committees are excellent at surfacing considerations and structurally incapable of being accountable, because accountability cannot be held jointly by people who can each reasonably assume another member is holding it.

2. Real authority to halt

That person can pause a deployment without prior consensus, and the organization has said so in writing. The authority is exercised rarely, and its value lies in being unambiguous before it is needed rather than negotiated during a crisis.

3. Clinical governance, not technology governance

Models are governed the way any other clinical intervention is governed. Local validation before go-live, monitoring afterward, defined withdrawal criteria, and integration into existing safety event reporting rather than a parallel process nobody uses.

4. A complete inventory, maintained

Including tools that arrived through procurement rather than through the electronic health record, and including the generative tooling that sits outside device authorization.² Most organizations discover on first attempt that they own more than they thought.

5. Clinicians consulted before adoption, not trained after it

Eighty five percent of physicians want to be consulted on or responsible for adoption in their practice.³ Organizations that involve them before selection get adoption. Organizations that involve them at go-live get override rates.¹³

6. Hiring for judgment, and paying for it

Treating this profile as a scarce strategic asset rather than a role to be filled. That means going to find people who are not looking, and it means recognizing that the individual capable of stopping a bad deployment is worth substantially more than the one who can explain a good one.

There is a reasonable objection to all of this, which is that it sounds like friction at a moment when boards want speed. The response is that the organizations moving fastest and most safely are the ones that resolved accountability early, because clarity about who decides is what removes the need to relitigate every deployment. Ambiguity is not speed. It is deferred cost.

PART XI

Conclusion: Who Signs

The public conversation continues to ask whether healthcare should use artificial intelligence. Three quarters of health systems, fourteen hundred authorized devices and four in five physicians have already answered.^{1,2,3} Continuing to debate the question is a way of avoiding the harder one.

The harder question is whether these organizations have leaders who can hold clinical judgment and computational judgment in the same head, who understand that a model validated elsewhere has not been validated here, who know that a strong accuracy figure can conceal a useless one, and who are prepared to put their name against the decision to keep a tool running or to switch it off. That profile barely existed five years ago, no career track produces it reliably, and demand for it is rising considerably faster than supply.

Meanwhile the consequences of getting it wrong currently land on the clinician at the bedside, because that is what the liability structure does with institutional ambiguity.⁴ Every organization that has not named an accountable executive has made that choice by default, and most have not realized they made it.

The systems that go looking for this leadership deliberately, and pay properly for it, will capture what the technology genuinely offers, and the benefit is real and already visible in the burnout data.¹⁰ The systems that fill the org chart and hope will find out what they bought at the worst possible moment, in a way that will be attributed afterward to a technology failure when it was a governance failure the whole time.

From the boardroom to the bedside, it comes down to a question every board should be able to answer before it is asked in a deposition. Who signs?

If the honest answer is the clinician who happened to be on shift, the organization has not deployed artificial intelligence. It has distributed liability and called it innovation.

REFERENCES

Sources

1. Eliciting Insights, 2026 AI Adoption Study, reported via Fierce Healthcare and Chief Healthcare Executive, 2026. Health system adoption, multi-solution deployment, and application-level adoption figures.
2. US Food and Drug Administration, Artificial Intelligence-Enabled Medical Device List, device counts as of August 2024, July 2025 and March 2026; annual authorization counts and specialty distribution as reported by MedTech Dive and The Medical Futurist. FDA final guidance on Predetermined Change Control Plans, August 2025.
3. American Medical Association, 2026 Physician Survey on Augmented Intelligence, March 2026; and AMA prior authorization survey, 2025. Adoption, trust, oversight, liability and training figures.
4. Analyses of AI and medical liability including Milbank Quarterly, "Artificial Intelligence and Liability in Medicine"; AMA Journal of Ethics, "Are Current Tort Liability Doctrines Adequate for Addressing Injury Caused by AI?"; and Johns Hopkins Carey Business School, "Fault lines in health care AI, part two." Nothing in this paper is legal advice.
5. Wong A, et al. External Validation of a Widely Implemented Proprietary Sepsis Prediction Model in Hospitalized Patients. *JAMA Internal Medicine*, 2021; and accompanying editorial, "The Epic Sepsis Model Falls Short, The Importance of External Validation."
6. Obermeyer Z, Powers B, Vogeli C, Mullainathan S. Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, October 2019.
7. Office of the National Coordinator for Health Information Technology, HTI-1 Final Rule, Decision Support Interventions certification criteria and predictive DSI source attributes, 2024.
8. The Joint Commission and Coalition for Health AI, Guidance on the Responsible Use of AI in Healthcare (RUAIH), September 17, 2025. Partnership announced June 2025.
9. The Joint Commission, voluntary Responsible Use of AI in Healthcare Certification, May 2026.
10. Ambient documentation literature, including a 2025 multi-site study of ambulatory clinicians across six health systems reported in *JAMA Network Open*; Mass General Brigham burnout reduction data; a randomized trial at UW Health; and a 2026 multi-center study of approximately 1,800 clinicians across five academic medical centers.
11. State legislative activity on artificial intelligence during 2025, as reported in healthcare governance analyses, 2026.
12. Chief AI Officer adoption and healthcare AI hiring trend data, 2026, including The Health Management Academy on AI governance leadership, NEJM AI on the Chief Health AI Officer role, and healthcare IT staffing search volume analyses.
13. Systematic review and meta-analysis literature on clinical decision support alert override rates and alert fatigue, including drug interaction alert override meta-analysis published in *Health Informatics Journal*, and alert fatigue measurement systematic review, 2026.

Lee Fossey leads healthcare executive search for ISG-Americas, advising health systems, behavioral health and post-acute platforms, and healthcare investors on confidential leadership search, from the boardroom to the bedside. This is No. 3 in a weekly series on the leadership questions shaping healthcare.

Figures are reported as published by the cited sources and were current at the time of writing, August 2026. Where studies disagree, both readings have been given. This paper is intended to inform leadership discussion and does not constitute legal, clinical or regulatory advice.