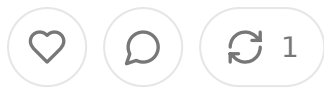


Nobody gets sued but the doctor: The legal vacuum at the center of the AI physician revolution

FEB 20, 2026 • PAID



Share

Abstract

Who absorbs liability when an AI-assisted clinical decision causes patient harm, what does the current legal ambiguity mean for founders, investors, and health systems deploying these tools at scale?

Key findings and data points:

- FDA has cleared over 1,300 AI-enabled medical devices as of late 2025, with 29 cleared in 2025 alone; 62% are software as a medical device (SaMD)
- 66% of U.S. physicians used AI in clinical practice in 2024, up from 38% in 2023. 78% single-year jump
- Malpractice claims involving AI tools increased 14% between 2022 and 2024; most involved diagnostic AI in radiology, cardiology, and oncology
- No existing federal law assigns liability to AI developers when an AI tool contributes to patient harm; current courts default to the treating physician
- The Federation of State Medical Boards recommended in April 2024 that clinicians, not vendors, be held liable for AI-generated errors
- GPT-4 outperformed physicians using GPT-4 in complex diagnostic cases in a study, suggesting the human-AI hybrid may actually underperform pure-AI in some contexts

- AI healthcare market estimated at \$39.25B in 2025, projected to reach \$504B by 2032 at ~44% CAGR
- AI-focused digital health deals in H1 2025 ran 83% larger than non-AI deals; \$3 of the \$6.4B raised went to AI companies
- Only 2% of U.S. radiology practices had integrated AI reading tools by 2024 despite hundreds of FDA-cleared options

The argument: The AI diagnostics revolution is outrunning the legal infrastructure designed to govern it. Physicians bear all the liability, vendors bear almost none, the incentive structures this creates are deeply misaligned with both patient safety and long-term enterprise value for companies in the space. The founders who figure this out early – and build accountability into their products rather than contract away – will have a significant structural advantage.

Table of Contents

The Setup: 950 Cleared Devices, 2% Adoption, and a Liability Cliff Nobody's Talking About

What the FDA Has and Hasn't Done

The Physician Gets the Bill: How Malpractice Law Currently Assigns Blame

Deskilling, Overdependence, and the Colonoscopy Problem

The Black Box Defense and Why It Won't Hold

What the EU Is Doing That the U.S. Isn't

Founder Implications: Build the Accountability Layer Now

Investment Thesis: The Liability Arbitrage Window Is Closing

The Setup: 950 Cleared Devices, 2% Adoption, and a Liability Cliff Nobody's Talking About

Here's a number that should make every health tech founder pause: as of late 2024, the FDA has cleared more than 1,300 AI-enabled medical devices. Radiology alone accounts for the overwhelming majority. And yet, a 2024 Associated Press survey found that only about 2% of U.S. radiology practices had actually integrated AI reading tools into their workflows. Think about that gap for a second. It's not a regulatory problem. It's not a technical problem. The tools exist, they're cleared, in many cases the clinical evidence behind them is legitimately strong. A large Swedish randomized trial found 17.6% higher cancer detection rates with AI-assisted mammography screening. Viz.ai's stroke detection algorithm hits AUC above 0.9 on retrospective datasets. Aidoc's intracranial hemorrhage tool reports sensitivity at 90% with low false-positive rates. The performance is real.

So what's the holdup? Some of it is workflow friction, integration headaches, and the usual institutional inertia that makes healthcare adoption timelines look like geological epochs. But a big, underappreciated chunk of it is liability terror. Clinicians are watching a legal landscape where they get to absorb all the downsides of an AI error, while the vendor takes the upside in the form of a recurring SaaS contract. That is a structurally bad deal, and the smarter physicians are figuring it out fast.

The broader picture is even more interesting. U.S. digital health startups raised over \$10B in H1 2025, with AI-focused companies capturing roughly \$3.95B of that, and those deals ran 83% larger than non-AI deals on average. Eight healthcare AI unicorns were minted in 2025. The money is flowing in fast and hard. Meanwhile, the legal infrastructure governing what happens when one of these tools contributes to a missed diagnosis or a treatment error is essentially non-existent in any coherent, codified sense. There is no federal statute. There is almost no AI-specific case law. There are strongly worded recommendations from medical boards telling physicians they're on the hook, and there are vendor contracts full of indemnification language.

that makes the physician's counsel wince. That combination – massive capital in real clinical deployment, and a liability vacuum – is the setup for a genuinely consequential legal and market reckoning.

What the FDA Has and Hasn't Done

Credit where it's due: the FDA has actually moved with more urgency on AI medical devices than most regulatory observers expected. In 2025 alone, 295 AI/ML medical devices received clearance, which is a record-breaking annual total. The agency cleared 223 in 2023 versus just six in 2015. That is a meaningful acceleration, and the Pre-Determined Change Control Plan (PCCP) pathway introduced in recent years is a genuine regulatory innovation – it allows AI developers to pre-specify how their algorithms will be updated and retrained without triggering a full new 510(k) submission every time, which was the bottleneck that was making continuous learning practically impossible under the old rules. That matters a lot for SaMD companies, which now represent 62% of all AI/ML clearances.

But here's what the FDA has not done: it has not created a liability framework. Its mandate is safety and efficacy. Whether a device is cleared tells you that the FDA found the evidence for its intended use sufficient. It says very little about what happens legally when the device is used appropriately and a patient is still harmed. The FDA cleared IDx-DR in 2018 as the first autonomous AI diagnostic device – screen for diabetic retinopathy without an ophthalmologist present, hitting 87% sensitivity and 90% specificity in a multicenter trial. But the FDA clearance letter doesn't tell you who gets sued if IDx-DR misses a case of disease in someone who loses vision. That question falls entirely into the tort system, which was not designed with autonomous AI in mind.

The January 2025 FDA draft guidance on AI lifecycle management is worth reading if you haven't. It proposes post-market monitoring requirements, performance drift detection, transparency standards around training data. All useful. None of it touches liability allocation. And importantly, 97% of AI devices have been cleared via the 510(k) pathway, which is the substantially equivalent to a predicate device route.

cross-sectional study of 903 devices published in JAMA Network Open found that at the time of regulatory approval, clinical performance studies were reported for only about 56% of analyzed devices. A quarter explicitly stated no performance study had been conducted. Less than a third provided sex-specific data. Under 25% addressed age-related subgroups. This is important context for anyone doing diligence on a diagnostics company: FDA clearance is a necessary but genuinely insufficient sign of real-world clinical robustness.

The Physician Gets the Bill: How Malpractice Law Currently Assigns Blame

Under existing U.S. malpractice doctrine, the standard is the reasonable physician acting under similar circumstances. It doesn't matter that an AI tool made a recommendation that any reasonable physician might have followed. If the outcome was bad, courts ask what the physician did, and the physician is typically the one who ends up liable. The Federation of State Medical Boards formalized this position in April 2024, explicitly recommending that its member boards hold clinicians, not AI makers, liable for AI-generated errors. The reasoning is that medical professionals are responsible for ensuring accuracy and veracity of evidence-based conclusions, regardless of what tool they used to get there.

The vicarious liability doctrine that some scholars have proposed – where an AI system acting as a clinical subordinate might have its errors attributed to the supervising physician or institution, similar to how an employer bears liability for employee negligence – remains largely theoretical in U.S. courts. There are no landmark AI malpractice cases yet. There have been claims involving diagnostic AI, and data suggests a 14% increase in malpractice claims touching AI tools between 2022 and 2024, mostly in radiology, cardiology, and oncology. But none have produced the kind of appellate precedent that would reshape how courts treat vendor liability.

What this means in practice is a wildly asymmetric risk structure. The vendor sells the tool, takes the subscription revenue, updates the algorithm (sometimes in ways the physician doesn't even know about via PCCP), and sits behind layers of contract

indemnification. The physician uses the tool in good faith, documents the encounter, makes a clinical judgment call that incorporates the AI output, and if anything goes wrong, they're the named defendant. This isn't hypothetical anymore. Malpractice insurers are responding in real time: some are now introducing AI-specific policy exclusions, others are requiring AI training as a condition of coverage. Jared Kaplan, CEO of Indigo, said publicly that while AI is a net positive that he expects will ultimately drive down malpractice rates, it exposes a new threat where there's no black-and-white answer about fault allocation. That's a pretty measured assessment. The practicing physicians navigating this daily would use less polished language.

The "no-win scenario" described by researchers at UT Austin and Johns Hopkins is worth unpacking because it maps directly to enterprise sales dynamics. Physicians are expected to critically evaluate AI outputs and override them when clinical judgment says to. But they're deploying these tools precisely because they're faster and often more accurate than human judgment alone. If you consistently override the AI, you're defeating the purpose of having it, and you might actually be making worse decisions. If you don't override it when you should, you're liable for blindly following an algorithm. Asking physicians to perfectly understand and apply AI while bearing full accountability is, as one researcher put it, like expecting pilots to design their own aircraft while flying. The analogy to aviation is interesting: in that industry, liability is distributed across pilots, systems, manufacturers, and regulators when automation fails. Healthcare has not made that structural shift.

Deskilling, Overdependence, and the Colonoscopy Problem

There's a subtler long-term issue that doesn't show up in malpractice statistics yet, and it's relevant to anyone underwriting AI clinical tools from an investor's standpoint. The ACCEPT trial in Poland is the clearest demonstration of it.

Endoscopists using an AI tool for polyp detection saw their adenoma detection rate improve meaningfully. When the AI was removed, those same endoscopists dropped from a 28% adenoma detection rate to 22%. Their underlying skill had atrophied.

because they'd been relying on the AI to do part of their cognitive work. This is deskilling, and it's not a fringe concern in the clinical literature anymore.

The colonoscopy case is actually relatively low stakes as an illustration. Scale this to breast radiology, where an AI second reader is now effectively standard of care in breast screening programs across parts of Europe, and the deskilling concern gets more serious. A 2024 survey of 572 European radiologists found widespread belief that AI was affecting their professional skills. And unlike the colonoscopy case, if a radiologist becomes overdependent on an AI screening tool and that tool has a systematic blind spot for a particular cancer morphology in a specific demographic, the deskilling effect and the algorithmic bias combine in a way that's very difficult to detect until you're looking at retrospective outcome data over years.

This matters for liability in a specific way: if it becomes established in the clinical literature that use of a particular AI tool degrades the long-term diagnostic capability of the physicians using it, the vendor starts to look a lot more like a defendant. The standard of care argument cuts both ways. Right now, courts don't expect clinicians to use AI. Soon, they may expect it. When that shift happens – and the trajectory suggests it's two to five years out in most high-volume diagnostic specialties – failure to use a cleared, well-evidenced AI tool might itself be considered a breach of the standard of care. At the same moment that becomes true, vendors who have been hiding behind physician liability will face product liability exposure they've never had to price into their contracts before.

The Black Box Defense and Why It Won't Hold

The opacity problem is real and underappreciated even among technical founders. Neural network-based diagnostic tools make predictions via mechanisms that are illegible to the physician, the hospital, or sometimes even the developers themselves. When a deep learning model trained on millions of chest X-rays flags a finding, it cannot produce a chain of reasoning the way a physician explaining their differential diagnosis can. This has obvious implications for malpractice defense: how does a

physician document their reasoning around an AI recommendation when they have access to why the AI made that recommendation? How does a plaintiff's attorney challenge an AI output in discovery when the training data is proprietary and the model architecture is a trade secret?

For the near term, this opacity actually protects vendors more than it hurts them; they can't easily prove the algorithm was negligent if you can't examine the algorithm. This is changing in multiple directions simultaneously. Explainable AI (XAI) is a hot and fast-developing technical subfield, with gradient-weighted class activation mapping and attention-based interpretability tools increasingly being embedded in clinical AI products to show which image regions drove a decision. The EU's AI Act, which has teeth and will apply to companies selling into European health systems, explicitly treats high-risk AI systems as requiring meaningful transparency and documentation of training data characteristics. Once explainability becomes standard for any company that wants European market access, it becomes the expectation everywhere, including in U.S. courtrooms.

The plaintiff bar is also getting smarter fast. Trial attorneys who litigated pharmaceutical product liability cases spent years learning how to depose biostatisticians and dissect clinical trial design. They will figure out AI, and they will hire the experts they need. The first case that successfully establishes vendor liability for a defective AI diagnostic algorithm – not the physician, not the hospital, but the company that built the model – will send shockwaves through the entire SaaS-in-healthcare pricing and contracting structure. Founders who are not thinking about what that case will look like, and what their product's documentation trail would show a plaintiff's expert, are building on sand.

What the EU Is Doing That the U.S. Isn't

The regulatory divergence between the U.S. and EU on AI liability is one of the most consequential dynamics in global health tech right now, and it's not getting near enough coverage. The EU AI Act classifies medical AI as high-risk and imposes requirements around transparency, human oversight, documentation, and post-r

monitoring that go meaningfully beyond FDA guidance. The EU AI Liability Directive goes further: it would apply non-fault liability rules to high-risk AI systems that cause harm, meaning plaintiffs would not need to prove negligence in the traditional sense. They'd need to show the AI was used, it was high-risk category, and harm resulted. That is a fundamentally different liability standard than anything currently codified in the U.S.

For U.S.-based founders, this creates a compliance forking problem. A company building an AI diagnostic tool that wants to sell into European health systems cannot build to FDA-minimum standards and expect that to be sufficient. They need to build to EU AI Act standards, which in practice means building to a higher document transparency, explainability, and oversight bar overall. Counterintuitively, this might actually be good for the market. Companies that build to EU standards will have products that are more defensible in U.S. litigation contexts as courts eventually demand more from vendors. The transparency and audit-trail practices required under the EU framework map well onto what U.S. plaintiffs' attorneys are going to want in discovery.

There's also a pure regulatory arbitrage angle here for investors: companies that are EU-compliant ahead of U.S. regulatory catch-up will have moats that are hard to replicate quickly once the U.S. does tighten standards. The companies that built robust privacy practices ahead of GDPR and CCPA had a real head start when those regulations rippled into procurement requirements. The same dynamic is coming with AI liability documentation.

Founder Implications: Build the Accountability Layer Now

The most actionable takeaway from all of this is simple in concept and hard in execution: the AI clinical tools companies that survive the first major liability wave will be the ones that built accountability infrastructure into the product, not the ones that contracted it away to physicians. What does that actually look like?

It starts with transparency documentation at the model level. Training data composition, demographic representation, out-of-distribution performance

characteristics, known failure modes – all of this should live somewhere that isn't confidential internal Confluence page. The JAMA Network Open study finding that less than a third of FDA-cleared AI devices reported sex-specific performance data and under a quarter reported age-specific subgroup performance is a roadmap for what plaintiffs will attack first. If a missed diagnosis involves a patient from a demographic that wasn't well-represented in training data, and that information wasn't disclosed to the deploying institution, the product liability exposure is significant.

It continues with contract structure. The current standard practice of pushing a liability via indemnification to the deploying health system is short-sighted. Health systems are increasingly sophisticated buyers who are seeing the same liability landscape and pushing back. Shared accountability structures – where the vendor accepts defined performance thresholds and contractual liability for outcomes where the model performs outside its validated operating envelope – are a competitive differentiator right now and will become table stakes. Companies like Aidoc and Viz.ai are starting to talk about outcome-based pricing structures. That's the direction this goes.

It extends to the clinician relationship. The deskilling problem is real and found should treat it as a product design challenge, not a physician education problem. Tools that surface their own uncertainty, that flag cases where they're operating at the edge of their training distribution, that prompt independent clinical review rather than passive acceptance, will have better outcomes data and better liability profiles. Building in structured override documentation – not just a checkbox but a lightweight clinician reasoning capture – creates the audit trail that makes both physician and the vendor defensible. This is good product design that happens to be good legal strategy.

Finally, invest seriously in post-market monitoring. The FDA's PCCP framework makes continuous model retraining possible. But continuous retraining without robust monitoring for performance drift and demographic bias creates exactly the kind of undetected systematic failure that turns into class action territory. A robust post-market performance program – with defined thresholds, demographic

stratification, and escalation protocols – is both a regulatory expectation under the new FDA draft guidance and a fundamental product safety requirement.

Investment Thesis: The Liability Arbitrage Window Is Closing

For angels and early-stage investors in health tech AI, the liability landscape creates a specific set of bets worth making. The first and most obvious is on companies building the accountability infrastructure layer itself. Clinical AI governance software, audit trail tooling, model performance monitoring platforms, explainability dashboards for clinical AI – these are picks-and-shovels plays that benefit from every dollar flowing into AI diagnostics regardless of which diagnostic models win. As health system procurement teams get more sophisticated about liability and as regulators tighten post-market requirements, the market for these tools grows without depending on any single clinical AI application succeeding.

The second bet is on companies with genuine EU regulatory compliance depth. Look for companies that have checked the box on CE marking, but companies that have actually built to EU AI Act standards around transparency, human oversight, and documentation. Those companies will be better positioned as U.S. standards tighten. They will have access to European health system customers that are increasingly moving faster on AI adoption than their American counterparts (partly because of clearer regulatory frameworks, ironically), and will be more defensible in litigation. The compliance depth is a real moat.

The third and more contrarian bet is on the outcome-based pricing pioneer. The AI diagnostics company that successfully executes an outcomes-linked contract with a major health system – with real financial accountability tied to defined performance metrics, not just an SLA – changes the entire enterprise sales conversation in the sector. It's a harder business model to build, it requires actuarial sophistication that most health tech companies don't have, and it requires confidence in your model's real-world generalizability that most companies also don't have. But the company

cracks it gets a defensibility story in enterprise procurement that no competitor match without taking on the same accountability.

The bet to avoid, or at least underwrite very carefully, is the company whose entire moat is model performance without a corresponding investment in the liability & accountability stack. In 2021, that was a fine company to build. In 2025, with malpractice claims involving AI up 14%, insurers writing exclusions, and the EU Liability Directive moving through ratification, the model-only company is accumulating tail risk that isn't priced into most cap tables. The best diagnostic accuracy in a market segment only generates durable enterprise value if the company survives the first serious adverse event and the litigation that follows it.

The AI physician revolution is real. The clinical evidence in a growing number of specialties is compelling. The market sizing is enormous and the capital markets confirming it. But the legal infrastructure governing how this all plays out is being built in real time, case by case, regulation by regulation. The founders and investors who engage seriously with that reality – not as a compliance burden but as a strategic variable – will have a substantial edge over those who don't.

[← Previous](#)

[Next →](#)

Discussion about this post

Comments

Restacks



Write a comment...

