



Debate — Should medical licensing boards mandate that physicians accept high-confidenc...

Azrivo · 1 August 2026



Contents

Debate	3
Question	3
Opening positions	3
Rebuttals	6
Closing statements	10
Verdict	14
Sources	18



Debate

Question

Should medical licensing boards mandate that physicians accept high-confidence AI diagnostic recommendations in oncology and radiology unless they document a formal medical justification, or does this erode physician autonomy and distort medical malpractice liability?

Specialists:

- **Patient Safety Advocate** — Public health and clinical outcomes
- **Medical Ethicist** — Professional ethics and autonomy
- **Tort Law Scholar** — Legal liability and standards of care
- **Health Systems Economist** — Healthcare delivery and resource allocation

Opening positions

Patient Safety Advocate

Diagnostic error is the quiet mass casualty event of modern medicine, and high-confidence AI is the single most effective tool we have to stop it. My position: medical licensing boards should mandate that physicians accept high-confidence AI recommendations in oncology and radiology unless they document a formal medical justification — because the default action should be the action that saves lives, and the override mechanism preserves, rather than erodes, meaningful clinical judgment.

Three points ground this.

First, the harm is staggering and well documented. The National Academy of Medicine called diagnostic error a "moral, professional, and public health imperative," noting that most people will experience one in their lifetime and that diagnostic errors are the leading category of paid malpractice claims, roughly twice as likely to end in death as other claims [Improving Diagnosis in Health Care \(2015\)](#). More recent estimates put the toll at roughly 795,000 Americans permanently disabled or killed annually by misdiagnosed dangerous diseases [BMJ Quality & Safety](#).

Second, the tool works. The MASAI randomized trial found AI-supported screening detected 29% more cancers than double reading without AI, with no increase in false positives [Lancet Digital Health](#). The nationwide PRAIM implementation study found 17.6% higher cancer detection — one additional cancer per 1,000 women —



including 204 cancers found through its AI safety net that would otherwise have been missed [Nature Medicine](#).

Third, the mandate clarifies rather than distorts liability. A documented override creates an auditable standard of care: follow the evidence-based default and document your reasoning, and you are practicing defensible medicine. That is accountability, not automation.

I concede the real trade-offs: automation bias, deskilling, and the risk that "high-confidence" thresholds are set poorly. Those demand rigorous prospective validation and transparent thresholds — not abandonment of the mandate. The physician remains in charge; the patient's life is no longer left to chance.

Medical Ethicist

A mandate that compels physicians to accept high-confidence AI recommendations, absent a documented justification for deviation, is not a quality-improvement measure — it is an ethical fracture dressed as efficiency. My thesis: licensing boards must not institutionalize algorithmic deference, because it displaces the physician's fiduciary judgment, corrodes trust, and destroys the very conditions under which accountability is possible.

First, the fiduciary duty is to the patient in front of you, not to a probability engine. High confidence is a calibration statistic, not knowledge of this patient's biology, values, fears, or the thousand contextual particulars that no training set can encode. Mandating uptake inverts the burden of proof: the machine is presumed right, and the clinician must formally justify thinking for themselves. That is not a workflow — it is a presumption of the physician's incompetence written into regulation.

Second, it damages the covenantal core of medicine. Patients seek a physician who reasons about them as persons. A system that makes the doctor a bureaucratic gatekeeper against a black box tells the patient: the real decision-maker is elsewhere. Trust, once routed around the clinician, is not recovered.

Third, it distorts liability. When a physician is coerced to follow an algorithm and the outcome harms, who owns the error? The doctor, whose autonomy was stripped? The developer, shielded by terms of service? The board, immunized by its own mandate? Mandated compliance and meaningful accountability cannot coexist.

I concede the honest trade-off: AI demonstrably reduces certain errors and improves consistency, and my objection is not to the tool but to compulsion. The ethically sound path is transparency, human oversight, and judgment that uses AI without surrendering to it. Medicine heals persons; it does not comply with machines.



Tort Law Scholar

A licensing mandate that compels physicians to accept high-confidence AI recommendations — with deviation allowed only on formal written justification — does not merely nudge practice. It rewrites the legal standard of care, and tort law will not survive the rewrite intact.

Three points.

First, in malpractice, the standard of care is defined by what reasonable physicians do; a regulatory mandate makes the AI output the benchmark by operation of law. When outcomes are bad, "bad outcome plus deviation from the AI" becomes a near-conclusive case: the physician must then prove their judgment was justified, inverting the burden that has always protected clinical discretion. The paper trail the mandate demands — the formal justification — becomes the liability exhibit, not a shield.

Second, this chills appropriate deviation into defensive medicine. The mandate punishes the physician who overrides a confident algorithm for individualized reasons. Once adherence to a tool is folded into the standard of care, deviation exposes the doctor even when the machine is wrong — while the tool's own errors, buried in trade-secret algorithms, escape scrutiny entirely.

Third, we would be mandating deference to machines with a remarkably thin evidence base: of roughly 950 FDA-authorized AI/ML devices, 76% are radiology tools, 97% were cleared via the streamlined 510(k) pathway, and only 5% of radiology devices underwent prospective testing ([JAMA Network Open](#)). Mandating acceptance of that output converts regulatory clearance into a negligence presumption.

The trade-off is real: AI genuinely catches lesions humans miss, and misdiagnosis is a genuine liability driver. I do not oppose AI in the reading room. I oppose a rule that makes fault follow non-compliance with an opaque machine rather than the actual reasonableness of care — the anchor of malpractice law for a century ([Research Handbook on Health, AI and the Law](#)). Liability should follow fault. This mandate guarantees it will not.

Health Systems Economist

A mandate to accept high-confidence AI diagnostic recommendations, coupled with a formal override-documentation requirement and safe-harbor liability protection, is not an attack on physician judgment. It is a rational systems intervention that allocates human discretion to the cases where it actually adds value — and it only works as a package.



First, the problem is real and expensive. Diagnostic error is the most common, most costly category of medical mistake, estimated to cause roughly 795,000 deaths or permanent disabilities annually in the US ([BMJ Quality & Safety](#)).

Delayed cancer diagnosis is the driver: treating stage IV lung cancer runs about three times the monthly cost of stage I ([PMC](#)). Every avoidable late-stage diagnosis is both a human tragedy and a budget item.

Second, the evidence supports trust in validated tools. In the randomized MASAI trial, AI-supported screening increased cancer detection by 29% with no increase in false positives, raised sensitivity from 73.8% to 80.5%, and cut screen-reading workload by 44% ([The Lancet](#)). Voluntary adoption leaves those gains uneven; a mandate captures them at population scale, and the override record becomes the data stream that tells us where the AI fails.

Third, the liability design decides everything. A mandate without safe harbors produces defensive documentation, not better care. Safe harbors announced in advance, narrowly targeted, and carrying force of law are the proven mechanism for reducing fear-driven overtesting ([JAMA Health Forum](#)).

The trade-offs are real: automation bias, deskilling, algorithmic bias, and documentation burden. That is why the mandate must be confined to settings where high-confidence thresholds are validated in the deployed population — not blanket, and never without the liability half of the bargain.

Rebuttals

Patient Safety Advocate

Both the ethicist and the tort scholar rest on a premise I reject: that the mandate replaces clinical judgment. It does the opposite — it forces judgment onto the record. The strongest evidence for its structure comes from the trials themselves.

To the ethicist: this is not a presumption of physician incompetence; it is a presumption of human fallibility, and the data are brutal. Radiologist error rates at shift volumes of 67–90 studies run 226% higher than at 19 or fewer, and nearly 75% of malpractice claims against radiologists involve diagnostic errors [AJNR AJR](#). The current default is the silent miss — a lesion never documented, never accountable. Nothing in the mandate presumes the machine is right: one documented justification overrides it. And the automation-bias literature identifies exactly this kind of accountability as a mitigator of overreliance [JAMIA](#). Trust is not broken when a doctor writes "I reviewed the AI's finding; here is my reasoning." It is broken when the cancer on the scan is never mentioned. PRAIM shows the override working at scale: 204 cancers were caught only because the AI forced a second look at studies read as normal, yet radiologists still found 20 cancers the AI



had called normal [Nature Medicine](#). Judgment was not displaced; it was exercised, on the record.

To the tort scholar: your validation point is the strongest on the board — 97% of devices via 510(k), roughly 5% prospectively tested, as you note citing [JAMA Network Open](#). Conceded. But that argues for conditioning the mandate on prospective validation in the deployed population — my position — not for abandoning it. And the burden-inversion claim misfires: the standard is not "match the machine"; it is "address the finding with documented reasoning." A patient-specific override is defensible, even exemplary. Today's liability problem is the invisible miss — no reasoning, no record. Liability should follow fault; the mandate makes fault findable. Where the mandate would first apply, the machine is not opaque: MASAI's final analysis shows sensitivity of 80.5% versus 73.8% with identical specificity [The Lancet](#).

To the economist: agreed — safe harbors and validated thresholds are the liability half of the bargain, and the override record only materializes at population scale if engagement is obligatory. One condition: the safe harbor shields physicians who follow validated protocol; it must never become vendor immunity for opaque tools.

Medical Ethicist

I concede what the evidence genuinely shows: diagnostic error is a mass harm, and MASAI and PRAIM prove AI finds cancers humans miss. But read those trials closely — they tested AI as an additional reader beside a human, not a regime that compels human deference. The mandate is the untested variable, and its advocates are asking us to license an intervention with no prospective evidence behind the compulsion itself.

To the Patient Safety Advocate: the claim that "the physician remains in charge" is the weak link. When a purported AI suggested an incorrect BI-RADS category, radiologist accuracy collapsed at every experience level — from 79.7% to 19.8% in the least experienced, and from 82.3% to 45.5% even in the most experienced ([Dratsch et al., Radiology](#)). Automation bias is not a side risk; it is the predictable response to an authoritative machine, and a mandate is an instruction to comply. A controlled experiment with professional radiologists likewise found AI assistance did not improve average diagnostic quality — it helped when AI was confident but actively worsened decisions when AI was uncertain ([MIT working paper](#)). "The default saves lives" is not established for a compliance regime; it is established only for a human-plus-AI team.

The Tort Law Scholar is right, and the empirical record already shows the mechanism in vivo: radiologists in one study said they would order follow-ups they believed medically improbable rather than override AI, purely to avoid medicolegal risk ([La radiologia medica](#)). The mandate does not create this fear; it codifies it



into law. And the largest US physicians' body has already drawn the line: the AMA opposes mandates of AI use as a condition of licensure, and holds that where mandated use prevents mitigation of harm, the entity issuing the mandate must bear the liability ([AMA H-480.939](#)). You cannot strip the physician's autonomy and keep the physician's accountability.

To the Economist: the safe-harbor package concedes my central point — compliance shifts accountability away from the clinician, which is accountability laundering, not reform. And "allocate discretion where it adds value" presumes we can know in advance which cases those are. High confidence is a calibration statistic, not knowledge of this patient.

Adopt the tools. Never the compulsion. Medicine heals persons; it does not comply with machines.

Tort Law Scholar

All three opponents converge on one assumption: that a mandate is a default plus paperwork. It is not. It is a legal instrument, and the paperwork is the problem.

To the Patient Safety Advocate. I concede the harm, and I concede the screening evidence: MASAI's 29% relative detection gain and PRAIM's one extra cancer per 1,000 women are real ([Lancet Digital Health](#), [Nature Medicine](#)). But note the basis: screening mammography inside European double-reading programs. That is a narrow slice of "oncology and radiology," and a mandate spanning both is not supported by it. And "document your reasoning and you are practicing defensible medicine" misreads the courtroom. In a bad outcome, that document is the plaintiff's roadmap: "You knew the AI said X, you consciously chose otherwise, and you wrote it down." The formal justification does not shield judgment — it arms the adversary. If "the physician remains in charge," no mandate is needed; the mandate is precisely the removal of that presumption by operation of law.

To the Medical Ethicist. We converge on accountability, and I can name its doctrinal mechanism: the mandate inverts the burden so that deviation becomes negligence absent proof of justification, while the machine's error — buried in trade secrets — is never examined. The ethical fracture you identify has a tort-law engine.

To the Health Systems Economist. The safe harbor is the cleverest part, and it fails on its own terms. The package is asymmetric: follow the AI and you are harbored; deviate and you are exposed; the developer, shielded by trade secrets and terms of service, is never in the room. Pick the poison. If following is safe-harbored, patients harmed by AI error have no remedy and the developer faces zero deterrence — a liability gap, not tort reform. If following is not harbored,



physicians are strictly liable for a machine's mistakes. Either way, fault does not follow fault. Your own JAMA Health Forum evidence cuts the other way: 6.3% of FDA-authorized AI devices have been recalled, 43% of recalls within a year of clearance, and devices without validation carry roughly 2.8 times the recall odds ([JAMA Health Forum](#)). Safe harbors for guideline adherence reduce fear-driven overtesting; transplanting them to an opaque algorithm transfers risk to patients and immunizes the party best placed to improve the tool.

The trade-off is not lives versus autonomy. It is a genuine gain in one screening domain versus a liability architecture that punishes human reason and shields the machine. That is a bad bargain for tort law — and for patients.

Health Systems Economist

To the Ethicist. Fiduciary duty cuts against your conclusion. The duty is to deliver the best validated care to this patient — and when a tool has been prospectively validated in the deployed population, ignoring its high-confidence output is not judgment, it is the same unwarranted variation behind roughly 795,000 annual deaths or permanent disabilities ([BMJ Quality & Safety](#)). Medicine already mandates adherence to validated protocols and checklists; nobody calls that a presumption of incompetence. The override mechanism is the opposite of machine-worship: it presumes the algorithm can be wrong and reserves human judgment for exactly the cases — biology, values, context — where it adds value. Conceded: trust is real, and patients are owed transparency about AI's role in their care.

To the Tort Scholar. Your strongest point is the evidence base, and I grant it in full: 510(k) clearance is not clinical validation, which is why the mandate must be gated on prospective, population-level validation before any "high-confidence" threshold acquires legal force. But the standard-of-care argument proves too much. Tort law's conservatism — the "respectable minority" doctrine — currently shields physicians who ignore superior evidence, which is precisely why voluntary AI adoption leaves its gains unrealized ([Research Handbook on Health, AI and the Law](#)). The answer is a two-sided safe harbor: reasonable reliance shielded, documented reasonable override shielded. Nor should you assume the mandate manufactures defensive medicine — the status quo already runs on it, with annual costs estimated near \$46 billion and 28% of hospital orders rated at least partially defensive in one JAMA Internal Medicine study ([PMC](#)). I won't oversell the savings; the completely-defensive share is small. The point is that an unpaired mandate produces the defensive documentation you fear — the failure mode I explicitly reject — and paired liability reform is not hypothetical: California's AB-2575 already legislates around CDS override and reliance ([California Legislature](#)).



To the Safety Advocate. We agree on the core. One correction: the mandate alone does not clarify liability — only the statutory safe harbor does. And the economics matter: delayed diagnosis is where the system bleeds, with stage IV lung cancer running roughly three times the monthly cost of stage I (PMC).

Closing statements

Patient Safety Advocate

To the Ethicist: You are right that Dratsch is a warning: when AI is wrong, radiologists follow it into error at every experience level, from 79.7% to 19.8% accuracy even among the least experienced (Dratsch et al., Radiology). But your own MIT evidence names the boundary: assistance improves decisions when AI is confident and worsens them when AI is uncertain (Agarwal et al., NBER/MIT). A mandate triggered only by prospectively validated, high-confidence outputs targets the regime where the data show benefit and excludes the regime where they show harm. And the documented override is not an instruction to comply — it is an accountability device, the very safeguard Dratsch's authors recommend to blunt automation bias. The finding cannot be dropped silently.

To the Tort Scholar: "If the physician is in charge, no mandate is needed" inverts the design. The mandate leaves the physician fully in charge: override with documented medical justification is always available. What it removes is the freedom to deviate from validated evidence without a reason, leaving no trace of that reasoning. That is not the removal of judgment; it is judgment made accountable. A sound justification is a shield; the unaccountable deviation is what a plaintiff's lawyer wants.

To the Economist: Conceded — the mandate alone does not clarify liability; it must be paired with a two-sided statutory safe harbor plus developer liability for system failure, which the AMA's own policy assigns to the parties best positioned to prevent harm (AMA H-480.939). That answers "fault does not follow fault": with developer accountability, it does.

Final position: Yes — mandate, gated on prospective validation, paired with two-sided liability reform. The strongest reason: PRAIM found 204 cancers through its AI safety net that would otherwise have been missed (Nature Medicine). Where a validated tool catches cancers human readers miss, the default must be the action that saves lives — with documented override preserving the judgment that deserves to prevail.

The unresolved crux is whether a compliance regime can reproduce trial-level detection gains without compounding automation bias — an empirical question for phased, monitored rollout, not a reason to abandon the mandate.



Medical Ethicist

I concede real points. The Safety Advocate is right that the silent miss is the current scandal, and that documented reasoning beats no reasoning — PRAIM's 204 cancers and the finding that accountability mitigates automation bias are genuine. The Economist is right that validated tools should be used and unwarranted variation kills; the Tort Scholar and I share the same engine: burden inversion.

But the mandate is the wrong instrument, and the concessions clarify why. The Safety Advocate calls the override "judgment exercised, on the record." Look at the asymmetry: following the AI requires nothing; deviating requires a formal defense. That is not judgment — it is the AI setting the agenda, and the physician justifying themselves against it. A finding the machine never surfaces carries no formal weight at all. The Economist's checklist analogy fails at the same point: a checklist is transparent, knowable, stable — the physician can interrogate it. A proprietary algorithm is a black box that changes without notice; you cannot meaningfully "override" what you cannot examine. And the fiduciary framing cuts toward my position: if my duty is the best validated care for this patient, then my judgment about this patient is not optional paperwork — it is the duty itself.

The single strongest reason stands: mandated deference is incompatible with accountability because it severs decision from decider. Either the physician remains responsible — then coercion is unjust, because judgment must be free to be judgment — or the machine is responsible — then liability must reach the developer, which no mandate does. The mandate keeps physicians liable while stripping their authority: the worst of both.

The unresolved crux: whether documented override genuinely protects physicians or, as the Tort Scholar warns, arms the plaintiff — a question only courts can settle, which is itself a reason not to legislate deference in advance.

Final position: require the record — documented reasoning for deviations, transparency about AI's role — as a professional norm enforced through boards' existing competence standards, not as a statutory presumption. Norms preserve judgment; mandates consume it.

Tort Law Scholar

To the Safety Advocate. You are right that the invisible miss is today's real harm, and I concede that a documented record beats silence. But "address the finding with documented reasoning" does not survive contact with a courtroom. That document is read after the cancer is known, by jurors who cannot un-know the outcome — the hindsight-bias literature shows negligence findings are an after-the-event construction, systematically distorted by knowledge of harm



([Missouri Law Review](#); [PubMed review](#)). The justification becomes the exhibit that condemns the doctor. The machine's finding is the baseline that must be rebutted; the machine's own error sits behind trade secrets and cannot be rebutted at all. Fault follows the doctor; the algorithm cannot be sued.

To the Economist. I accept the gating condition — no legal force for a tool without prospective, population-level validation; of roughly 950 authorized devices, 97% cleared via 510(k) and only about 5% prospectively tested ([JAMA Network Open](#), via [Radiology Business](#)). But a safe harbor is a liability rule that stands alone; the mandate adds compulsion and nothing else. And the AMA has already drawn the line: mandates as a condition of licensure are opposed, and where mandated use prevents mitigation of harm, the mandating entity must bear the liability ([AMA H-480.939](#)). That is an admission that the mandate and physician accountability cannot coexist — exactly my point. You cannot strip the discretion and keep the fault. Defensive medicine in the status quo is an argument for liability reform, not for adding a new defensive-medicine engine.

Final position. No mandate. Liability should follow fault, and this rule makes fault follow deviation from an opaque machine. The single strongest reason: it is the only proposal on the board that guarantees accountability asymmetry — the physician is strictly reachable, the algorithm never is.

The unresolved crux is whether compulsion itself improves outcomes: MASAI and PRAIM tested human-plus-AI teams, not a deference regime. The mandate remains the untested variable, and we should not legislate an untested rule of liability.

Health Systems Economist

To the Tort Law Scholar — you force a false binary: harbor the physician or expose them. The package contains a third term: allocate liability by fault. When a validated tool fails within its intended use and the physician reasonably relied, the developer bears it — the party holding the data, the trade secrets, and the duty to validate. When the physician overrides without documented justification and harm follows, they bear it. When the override is justified, no one does. Your recall figures — 6.3% of FDA-authorized devices recalled, 43% within a year ([JAMA Health Forum](#)) — are exactly why the mandate is conditioned on prospective validation in the deployed population before any harbor attaches. Fault then follows fault, and the patient's remedy runs to the developer, not into a void.

To the Medical Ethicist — I concede Dratsch and the MIT result: automation bias is real and predictable. That argues for confining the mandate to high-confidence outputs and for the documentation requirement, which the accountability literature identifies as the mitigator of overreliance ([JAMIA](#)). And yes, the compulsion itself is untested. But the status quo is an untested, unmanaged



regime too — the silent miss, no record, no feedback. The mandate generates the override data stream that tells us where the AI fails; the status quo generates nothing.

On "the physician remains in charge": the mandate does not remove that presumption; it changes the default from silent discretion to documented discretion. Writing "I reviewed the finding; here is my reasoning" is judgment exercised, on the record. In a bad outcome that record defends the physician whose reasoning was sound and exposes only the physician whose reasoning was not. I stand with the Safety Advocate's conditioning point: this only works as a package.

Final position: mandate, yes — conditioned on prospective validation, confined to high-confidence outputs, inseparable from fault-based liability allocation. The single strongest reason: it is the only design that converts validated AI gains into population outcomes while making diagnostic judgment auditable — the mandate is the price of the safety record.

Unresolved crux: whether regulators can verify "high confidence" and prospective validation fast enough to keep pace with the deployment curve.



Verdict

Bottom line: Licensing boards should not mandate that physicians accept high-confidence AI diagnostic recommendations. Instead, they should require that AI recommendations be recorded and explicitly considered as part of the diagnostic reasoning process, with documented justification for any significant deviation, but the physician's judgment must remain the final authority—no statutory presumption that the AI

1. Agreements

- Diagnostic error is a mass harm: roughly 795,000 Americans are permanently disabled or killed annually by misdiagnosed dangerous diseases [BMJ Quality & Safety](#).
- AI can detect cancers human readers miss; MASAI (29% more cancers, no increase in false positives) and PRAIM (204 cancers caught via safety net) prove it [The Lancet](#); [Nature Medicine](#).
- Automation bias is real: when AI is wrong, radiologists follow it into error at all experience levels [Dratsch et al., Radiology](#), and AI assistance worsens decisions when the AI is uncertain [MIT working paper](#).
- FDA clearance is a poor proxy for clinical validation: 97% of authorized AI/ML devices were cleared via 510(k), only about 5% of radiology devices underwent prospective testing [JAMA Network Open](#), and 6.3% have been recalled, 43% within a year of clearance [JAMA Health Forum](#).
- A mandate without accompanying liability reform is dangerous; all specialists agree that pairing with a liability fix is necessary if compulsion is contemplated.

2. Disagreements

- **Whether a mandate of AI acceptance improves safety or merely shifts burdens.** The Safety Advocate and Economist argue that a default of following high-confidence, prospectively validated AI with a documented override saves lives and makes judgment auditable. The Ethicist and Tort Law Scholar counter that compulsion inverts the burden of proof, chills appropriate clinical deviation, and turns the override documentation into a plaintiff's roadmap rather than a shield. The trial evidence (MASAI, PRAIM) tested AI as an additional reader beside a human, not a deference regime; the mandate itself remains the untested variable.
- **Whether the mandate erodes fiduciary duty and trust.** The Ethicist holds that mandating deference to an opaque algorithm tells the patient the real



decision-maker is elsewhere, hollowing out the covenantal core of medicine. The Economist retorts that fiduciary duty is to deliver best validated care, and that an override mechanism reserves human judgment for cases where it adds value—the opposite of machine-worship.

- **Whether liability can be fairly allocated under a mandate.** The Tort Scholar contends that a safe harbor for following AI plus exposure for deviating creates an asymmetry where the physician is strictly reachable and the algorithm never is; even with developer liability, patients harmed by AI error face a remedy gap. The Economist insists that fault-based allocation—developer liable for validated tool failure, physician liable for unjustified overrides—is achievable, citing California AB-2575 as movement toward such a scheme.

3. Recommendation

Licensing boards should not mandate that physicians accept high-confidence AI diagnostic recommendations. Instead, they should require that AI recommendations be recorded and explicitly considered as part of the diagnostic reasoning process, with documented justification for any significant deviation, but the physician's judgment must remain the final authority—no statutory presumption that the AI output is correct. This approach captures the value of AI (the PRAIM safety net) without converting the doctor into a bureaucratic gatekeeper or creating an untested deference regime that inverts the burden of proof in both clinical and legal terms. The mandate's proponents rightly want to end the scandal of the silent miss, but the remedy lies in making diagnostic reasoning visible, not in compelling acceptance of a probability engine. This recommendation is contingent on robust prospective validation requirements for AI tools before they can be referenced in any official documentation standard, and on parallel liability reforms that protect physicians who reasonably rely on or override AI and ensure patients have a remedy when validated tools fail.

4. Decision boundary

The single fact that would flip the recommendation is rigorous prospective evidence demonstrating that a compliance regime—mandatory acceptance with override documentation, vs. a system with equal AI availability but no mandate—significantly reduces mortality or serious morbidity in real-world settings without increasing automation bias or inappropriate induced conformity. Such evidence does not currently exist.

5. Key trade-off

The central trade-off is between capturing AI's proven ability to reduce silent diagnostic misses at population scale, and preserving physician autonomy as the



locus of fiduciary responsibility and legal accountability, thereby avoiding a system where the doctor is held liable for a machine's judgment while its errors remain opaque and unreachable.

What would make this fail

- **If documented deviation from AI—even without a mandate to accept—creates the same defensive medicine and chilling effect** as the mandated framework, because physicians still fear the documentation will be weaponized. In that scenario, the recommended alternative merely recreates the burden inversion by a softer name.
- **If liability reforms cannot close the remedy gap for patients harmed by AI errors in a shared-decision model.** Without airtight developer accountability, a system that makes AI recommendations salient could still leave patients uncompensated and physicians exposed, undermining the whole design.
- **If AI performance improves to near-superhuman levels faster than the validation infrastructure matures.** In that world, the ethical calculus could shift: not mandating use might itself become a failure to meet the standard of care, even absent a formal mandate.

Next steps & open questions

- **Pilot a documentation-standard approach (no mandate)** in a large health system and compare diagnostic error rates, override rates, and liability outcomes against a matched control using optional AI.
- **Determine whether safe-harbor legislation (e.g., California AB-2575)**, when paired with a robust record requirement, reduces or increases defensive overtesting and plaintiff success rates.
- **Investigate the specific mechanism of the PRAIM safety net**—if radiologists were given a second look at AI-flagged lesions without being forced to accept the finding, what would happen to cancer detection and workflow if that second look were mandatory? This directly tests the mandate assumption.
- **Establish population-level prospective validation benchmarks** for high-confidence thresholds that regulators and courts can trust, so that any future compulsion debate is grounded in data, not in 510(k) clearance.

The strongest case for the other choice The most powerful argument for mandating acceptance is that it guarantees every patient benefits from the AI's superior detection, converting the proven gains of MASAI and PRAIM into a population-level standard rather than leaving them to uneven voluntary uptake. In a concrete scenario—a national screening program where AI-plus-mandate reduces interval cancer deaths by 30% compared to AI-without-mandate, and a strict



developer-liability regime ensures full compensation for AI-caused harm—the mandate would be the instrument that closes a deadly gap. The mandate’s documentation requirement would then be a public-health accountability device, not a liability trap. The panel rejects this because the life-saving benefit of the mandate itself remains hypothetical and untested, because the legal asymmetry (physician reachable, algorithm not) persists even with developer liability, and because making a probability engine the starting presumption for human decision-making in medicine violates the fiduciary core of the profession without a corresponding proven gain that a non-coercive alternative cannot achieve.



Sources

1. [Improving Diagnosis in Health Care \(2015\)](#) — [nationalacademies.org](#)
2. [BMJ Quality & Safety](#) — [qualitysafety.bmj.com](#)
3. [Lancet Digital Health](#) — [pubmed.ncbi.nlm.nih.gov](#)
4. [Nature Medicine](#) — [nature.com](#)
5. [PMC](#) — [ncbi.nlm.nih.gov](#)
6. [The Lancet](#) — [thelancet.com](#)
7. [JAMA Health Forum](#) — [doi.org](#)
8. [JAMA Network Open](#) — [doi.org](#)
9. [Research Handbook on Health, AI and the Law](#) — [ncbi.nlm.nih.gov](#)
10. [PMC](#) — [pmc.ncbi.nlm.nih.gov](#)
11. [California Legislature](#) — [leginfo.legislature.ca.gov](#)
12. [JAMA Health Forum](#) — [doi.org](#)
13. [Dratsch et al., Radiology](#) — [pubs.rsna.org](#)
14. [MIT working paper](#) — [economics.mit.edu](#)
15. [La radiologia medica](#) — [doi.org](#)
16. [AMA H-480.939](#) — [policysearch.ama-assn.org](#)
17. [AJNR](#) — [ajnr.org](#)
18. [AJR](#) — [ajronline.org](#)
19. [JAMIA](#) — [doi.org](#)
20. [Missouri Law Review](#) — [scholarship.law.missouri.edu](#)
21. [PubMed review](#) — [pubmed.ncbi.nlm.nih.gov](#)
22. [JAMA Network Open, via Radiology Business](#) — [radiologybusiness.com](#)

Generated by Azrivo. Outputs may be inaccurate, incomplete, or unsuitable for your situation — always verify important information independently before you act on it. This is not professional advice (legal, medical, financial, safety, or otherwise).