

Same Benchmark Performance, Opposite Care

Validation shows a system could be used. Governance is what establishes that it may keep being used

In June 2026, a study in Nature Medicine reported that three general-purpose frontier language models, GPT-5.2, Gemini 3.1 Pro and Claude Opus 4.6, outperformed two specialized clinical AI tools, OpenEvidence and UpToDate Expert AI, across medical benchmarks and a set of 100 real physician queries collected from a single live clinical deployment (Vishwanath et al., Nature Medicine, 12 June 2026). The headline is powerful, which makes it dangerous, and it should be read with the discipline the rest of this article defends: before it means anything, one has to know the tasks, the systems, the models, the endpoint, the definition of "real-world," and whether the intended uses were even commensurable. The tools compared were commercial specialized systems, not devices carrying a regulatory clearance, so the popular gloss that "general-purpose AI beats clinical AI" already smuggles in one inferential jump this article exists to resist.

The field made the point itself within the quarter: on 3 September 2026 Nature Medicine published a Matters Arising (Beaulieu-Jones and Nemati) arguing that the limited benchmarks constrain the study's conclusions, precisely when they are stretched toward procurement, reimbursement, or regulation, three different decisions taken by three different authorities. The study authors replied the same month, maintaining their conclusion under their own deployment conditions, which is itself the point: the disagreement is not about the numbers, it is about how far a benchmark travels toward a decision. Take the finding narrowly, then. It does not show that clinical AI fails. It shows that systems ranked identically on a benchmark can sit very differently once a real decision, and a real decider, are in the room. That two systems with the same benchmark performance can produce opposite care is the door. The interesting question is not why. It is what this tells us about how AI should be governed.

Here is the answer in one line, before the argument that earns it: validation establishes that a system could be used at the moment it was assessed; governance must establish that it may keep being used. A clearance does not certify a number, it accepts a passage from an observed performance to a claim, and from a claim to a decision, and that passage has to be maintained, not merely granted. Three things must be kept apart, because the discourse on AI collapses them and pays for it. Epistemic credibility is a property of a claim against its evidence: what may we reasonably believe. Decisional admissibility is not a property of the model at all, it is a relation: a claim is never

admissible in the abstract, only for a specified decision, under a specified authority, loss, and consequence. Institutional authorization is a third thing again, the act by which a regulator, a payer, or a clinical service permits use. Three objects, three authorities, three standards of proof. Performance does not imply validity, validity does not imply utility, utility does not imply a decision, and none of them, established once, implies admissibility later. The last relation is the one this article is about, and it is the one neither a benchmark nor an initial validation can reach.

Four things a score is not

Separate, bluntly, four claims that the word "performance" runs together. Discrimination is not calibration. Calibration is not clinical utility. Clinical utility is not patient benefit. A model can be excellent at the first and worthless at the last, and each step needs its own justification. Most disputes about clinical AI are disputes about which rung someone is standing on without saying so.

The Nature Medicine study illustrates the compression directly. Its headline verb, "outperform," bundles four graded dimensions the clinicians actually scored apart, correctness, completeness, safety, and clarity, and the systems did not separate the same way on each: the frontier models pulled ahead more on clarity than on clinical correctness, a Kendall's W of 0.292 against 0.141, and no significant difference was found on harmful content or hallucination at all. "Better," here, is largely "clearer," which is not nothing and is not the same claim. One tool also declined nearly a fifth of the queries where the frontier models declined almost none, so a refusal policy, a design choice, is silently folded into the number that reads as performance. A score is rarely a measure of one thing.

Begin then with the estimate, since it settles the matter with no appeal to distribution shift. A point estimate is not a posterior: the same 0.87 can sit in a tight interval or a wide one, and the uncertainty is not of one kind. Sampling uncertainty can often be reduced by more representative observations; measurement, transport, and structural uncertainty generally require different evidence, not merely more of it. It is that second family a single figure hides most completely, and in the high-dimension, low-sample regime common in clinical machine learning, that distinction becomes consequential.

Our ISPOR work in rare BRAF V600E metastatic NSCLC provides a useful illustration. From a real-world cohort of only 184 patients, Gaussian Copula generation achieved an aggregate fidelity score of 0.954 (bootstrap 95% CI 0.941 to 0.957) and 99.7% machine-learning utility. Yet those reassuring global figures did not imply identity at the level of clinically meaningful estimates: 12-month overall survival differed by 8.9 percentage points, and the hazard ratio associated with ECOG performance status of 2 or more was 2.75 in the real cohort against 1.21 in the synthetic one. The direction of the prognostic effects stayed stable across 100 synthetic replications, but their magnitude did not. A

high aggregate score can therefore coexist with substantial uncertainty at precisely the level at which a clinical decision may depend on the estimate.

Calibration makes the point canonical. Take two models with identical benchmark discrimination, calibrated in the evaluation population. For one patient, the first returns a probability of 0.32 and the second 0.18. At a therapeutic threshold of 0.25, the first says treat and the second says withhold. Two senses of "score" are in play, and the title turns on their difference: the aggregate metric is the same, the individual prediction is not. Same benchmark performance, opposite care, for the same patient. Calibration is also distribution-dependent, so it is established in internal validity and must be preserved in transport, which is already a first sign that a property proven once is not owned forever.

One further distinction, the one most often skipped. A predicted risk is not a treatment effect: the patient at highest risk is not necessarily the one who gains most from acting. This is where credibility and admissibility part company. The claim "this patient has a 38 percent risk of progression" can be entirely credible and still say nothing about whether to treat. Credibility is earned against evidence; admissibility is earned against a decision.

What turns evidence into a decision, and what can defeat it

If a score is not the evidence, what is? The passage from data to decision runs through operations, measurement, identification, estimation, transport, action, and the useful word for what licenses each step is warrant, in Toulmin's sense. A warrant is simply the reason evidence is allowed to support a claim: the evidence, plus the assumptions required to make the inference. Evidence and assumptions are not the same thing, which matters in an article about the burden of assumptions: evidence supports a claim, an assumption licenses an inference the evidence alone does not carry.

Three warrants span the chain. The internal warrant takes measurement to a trustworthy estimate: construct, sound measurement, identification, estimation uncertainty, calibration in the evaluation population. It can fail with no shift at all, through a mismeasured endpoint, a data leak, uncorrected multiplicity, or a cohort too small to calibrate. The transport warrant takes the estimate to the target population and system, calibration preservation included. The actionability warrant, usually left unstated, takes a transported estimate to an action: it establishes that using the estimate changes the preferred action under the relevant decision rule, utility, constraints, and decision-maker. The individual patient, for whom acting must beat not acting, is the simplest case; a payer or a health system optimizes under scarcity and population benefit.

Each critical assumption should be interrogated for how it can be contested, and not every assumption can be contested the same way. Some carry an observable defeater, a

finding that would directly invalidate them, and can become a monitoring signal. Some admit only an indirect challenge, some a sensitivity test, some external evidence, and some no empirical test at all, only argument. The epistemic fragility of a claim can be read from the share of its architecture that rests on assumptions which are weakly testable, and pretending that all epistemic uncertainty converts into a dashboard metric would be a new version of the illusion this article resists. A validation study, then, does not stand against justification, it is part of it: a prospective, external study contributes evidence, it never abolishes the transport assumptions, it makes them more or less defensible. The audit of fifty-three medical benchmarks at ACL 2026 shows how routinely the internal warrant fails on its own terms before transport is even in question.

Not a checklist but an argument, and why the same evidence licenses opposite decisions

The honest way to present this is as a unification, not an invention, since three traditions already hold parts of the ground. The Context of Use, central to intended-use doctrine and to the FDA's risk-based credibility framework for drug and biologics decisions, specifies the target domain. Decision Curve Analysis folds in the loss and the threshold and turns performance into net benefit (Vickers and Elkin). External validation addresses distributional shift. None, alone, delivers decisional sufficiency. The Context of Use says where the performance must hold; the justification says why the observed performance is credible evidence that it will. Decision Curve Analysis is a bridge, not a rival: it computes net benefit after assuming its probabilities are valid and transportable, so it arrives after the operation it presupposes. Net benefit is not transportability.

What the architecture adds is a threshold whose shape matters. A justification need not be true, only credible enough for the consequence, and the credibility of a set of assumptions is not a scalar to be averaged, because some are non-compensatory: a false critical assumption is not redeemed by five well-supported ones. Nor are assumptions independent; some are substitutable, some complementary, some critical only conditionally on another. The architecture is therefore not a checklist but an argument, best read as a graph with two kinds of edge, those that support a claim and those that attack it, since a defeater stands against a warrant. This flips the posture of monitoring. Governance does not only accumulate evidence that the system works; it searches for evidence capable of defeating the current justification. The question moves from "show that it still works" to "look for the reason the authorization should no longer hold," and reading the graph is what reveals the single points of epistemic failure, the assumptions on which too many claims silently depend.

This is also why the same evidence can license opposite decisions, and the general form matters more than the clinical slogan. A decision is a function of the evidence, the utility,

the constraints, and the authority, not of the evidence alone. Identical epistemic evidence with different loss functions is not the same decision problem, so two rational actors can choose opposite actions from one number. The regulator, the payer, and the clinician are the standing example and they do not decide the same thing: the regulator sets availability, the payer sets accessibility, the clinician chooses care. The Matters Arising named the same trio, warning against carrying one benchmark toward procurement, reimbursement, and regulation as if they were one question. Same evidence, opposite decisions, each admissible for its own decision, authority, and loss. This generalizes at once beyond medicine, to a credit model read by a risk officer and a regulator, a screening tool read by a recruiter and a court, a detector read by an analyst and an auditor. Opposite outcomes from one model are not always error; often they are correct answers to different questions asked of the same estimate.

Transport contracts, authority, and governance in time

Transport needs an object precise enough to be falsifiable, or it dissolves into the truism that context matters. That object is the invariant, written as a contract in the engineering sense, an explicit set of conditions under which a claim remains admissible, not a legal instrument. Transport becomes operational by naming a minimal set of properties whose preservation is necessary, under the causal and measurement assumptions of the claim, for the evidence to stay relevant in the target setting. Necessary, not sufficient: even if the named invariants hold, an omitted variable can break the claim, which is why this is engineering and not a theorem. The contract states the property, the tolerance, the evidence that it holds, the monitor that watches it, the authority that owns the decision to act, and the cadence at which the evidence itself expires. A crossed threshold, moreover, is not a decision but the start of a graded response, alert, investigation, restricted use, revalidation, suspension, each step owned by someone. This is the load-bearing point for anyone who has to run such a system: a monitoring threshold without an authority and an action policy is telemetry, not governance. And two systems equally justified today can differ sharply in what it costs to keep them justified, one resting on observable, quickly testable assumptions, the other on structural assumptions revisable only by external study; the second carries more epistemic debt, and its lifecycle cost is higher for the same claim.

The regulatory frame illustrates the multi-authority point, and it is an illustration, not a foundation, since the doctrine holds whatever the calendar does. The stable, structural fact is this: where an AI system is, or is a safety component of, a medical device requiring notified-body assessment, its AI Act obligations are folded into the existing MDR or IVDR conformity assessment rather than run as a separate procedure, and several authorities then read the same evidence for different losses. The timetable is contingent: general high-risk obligations were set for August 2026, while for AI embedded in regulated

medical devices the Council agreed in March 2026 to extend the deadline further, toward 2028. The joint FDA and EMA Guiding Principles of Good AI Practice in Drug Development, issued in January 2026, are a half-right ally here, naming lifecycle monitoring and a risk-based approach while stopping short of an inspectable, owned, defeasible justification.

Which yields the definition worth keeping. AI governance is the controlled maintenance of decision admissibility under changing evidence, assumptions, context, and consequences. Its function is not conservative. A good governance system does not always keep a warrant alive; sometimes its duty is to end it. It does not keep warrants alive, it keeps them defeasible.

Coda: the harder problem, and what a governed decision means

One case sits beyond even this, and it is the subject of a companion article, so a sentence must suffice here. The doctrine so far assumes monitoring means watching whether the world changed around the model. There is a harder case: a system that acts on the process it observes changes that process, prediction leading to action, action to outcome, outcome to the next dataset. Once prediction changes treatment, post-market monitoring of intervention-coupled AI becomes a causal-inference problem, not merely a performance-monitoring one, and the same is true, more sharply, when the decisive quantity is a counterfactual that is never observed, as with the synthetic control arms and digital twins now entering regulatory evidence. That is where the burden of decision-relevant untestable assumptions is greatest, and where this argument continues next.

Which returns us to Elsa, as illustration and not as proof. The regulator that deploys a generative tool prone to inventing studies meets, in its own instruments, the confusion it must police in the ones it evaluates: an output that reads as evidence is not thereby evidence, and fluency is not a warrant. Elsa does not validate the doctrine, it stages it. Models produce estimates. Warrants make claims credible. Decisions make them consequential. Governance keeps the justification defeasible. Validation asks whether the system was justified at the moment it was assessed. Governance asks whether we can still justify using it now, who holds the authority to decide otherwise, and what evidence would force us to stop. A governed decision is not one whose model was validated, but one whose justification stays contestable over time, which is also why, to the mild disappointment of those who like dashboards, it will not reduce to a single governance score.

Bibliography

- Vishwanath K, Alyakin A, Ghosh M, et al. General-purpose large language models outperform specialized clinical AI tools on medical benchmarks. *Nature Medicine*. 2026;32:2405-2409. Online 12 June 2026. DOI: 10.1038/s41591-026-04431-5.
- Beaulieu-Jones B, Nemati S. Limited benchmarks constrain the conclusions of a general-purpose versus clinical AI comparison (Matters Arising). *Nature Medicine*. Published 3 September 2026. DOI: 10.1038/s41591-026-04638-6.
- Vishwanath K, Alyakin A, Ghosh M, et al. Reply to: Limited benchmarks constrain the conclusions of a general-purpose versus clinical AI comparison. *Nature Medicine*. Published September 2026. DOI: 10.1038/s41591-026-04637-7.
- "Beyond the Leaderboard: Rethinking Medical Benchmarks for Large Language Models." ArXiv 2508.04325 (ACL 2026).
- Vickers AJ, Elkin EB. Decision curve analysis: a novel method for evaluating prediction models. *Medical Decision Making*. 2006;26(6):565-574.
- FDA and EMA. Guiding Principles of Good AI Practice in Drug Development. 14 January 2026.
- Tadmouri A, Barkaoui S, Bennani M, Vetillard J, Mejbri H. Assessing Reproducibility of Real-World Evidence Analyses Using Synthetic Data: A Validation Study in Rare BRAF V600E Metastatic Non-Small Cell Lung Cancer. ISPOR Europe 2026 (Top 5% abstracts).