

WINNING EITHER WAY

An open playbook for scaling clinical AI in medtech, 2026-2030

Author: Bogdan Rynkowski

Date: July 2026

Audience: Innovation, strategy, and digital leaders in clinical medtech, and the consultancies that serve them
*Developed on the Oxford AI-Driven Business Transformation Executive Programme, Saïd Business School, University of Oxford, and
matured across 25 years in health technology.*

Version 1.0, July 2026. Signal states and enforcement dates reflect the regulatory position as of this date.

Summary

Clinical AI confronts medtech executives with a structural strategic problem: deep uncertainty about whether clinical and regulatory trust in AI holds or fractures by 2030. This paper applies the Oxford Scenarios Programme methodology to that uncertainty, using two futures to bracket it. In Frozen Trust, a run of clinical AI failures triggers moratoriums and procurement withdrawal, and power concentrates around regulators and liability insurers. In Full Commitment, accumulated clinical evidence drives platform-led consolidation, and power migrates to the integration layer and outcome data. Strategic plays across Purposes, Players, Partnerships, and Processes are developed for each future and tested for robustness across both. Five no-regret moves perform in either one, anchored by a jurisdiction-portable change-control architecture aligned to FDA, EU AI Act, and NMPA requirements: a regulatory moat under Frozen Trust, a scaling engine under Full Commitment. The paper closes with the one proposal worth sizing first, whether clinical AI validation infrastructure can transition from a governance requirement into a recurring-revenue growth engine, and with a note on who wrote it and why.

1. Purpose

A doctor looks at an MRI scan. Her training is excellent. What she lacks is the pattern recognition of ten thousand similar cases, the kind that takes decades to build and cannot travel with her to the next understaffed clinic. Clinical AI can put that within reach now. Which company gets it to her, governed and at scale, is the open question this paper addresses.

The paper is about the one uncertainty that decides how: whether clinical and regulatory trust in AI holds or fractures by 2030. It does not predict which future arrives. It identifies the moves that pay off in either one. The argument is written as a narrative rather than a deck because the reasoning has to hold on the page. Every claim is either evidenced or named as a bet.

Industry survey data sets the starting point. AI already works in clinical settings: half of clinicians report that AI has increased their capacity to see more patients, with a median gain of eight patients per week among them, and the question has shifted from whether to scale to how (Partovi, 2026, drawing on the 2026 Future Health Index). The constraint that shapes the answer is capital discipline. After years of margin pressure, most incumbent balance sheets leave little room for new platform capex. The plays that win must generate recurring revenue from assets a company already owns.

2. Tenets

Five tenets govern the recommendations. First, legitimacy is the scarce input in clinical AI; capability is abundant and getting cheaper. Second, AI is only real if it produces measurable clinical gains; anything else is theater. Third, governance built once to a shared standard is an asset that travels across jurisdictions, while governance rebuilt per market is a cost that compounds. Fourth, humans belong above the loop, setting the rules the machine enforces, not rubber-stamping outputs they cannot trace. Fifth, a move that pays in both futures costs no more to make than a move that pays in one; commit to those first.

3. Where the market stands in 2026

Every major medtech player is building AI for productivity. Arguably none has yet fielded a proven, stable, governed AI for clinical decision support at scale. That gap is the opening, and closing it is execution work. The window is open. Not for long.

Three regulators are converging independently on the same architecture. The FDA's Total Product Life Cycle framework, the EU AI Act's Article 6 high-risk classification, and China's NMPA YY/T 1406-2026 standard all now expect continuous lifecycle risk management instead of point-in-time approval, transparency built into submission and labelling, and closed-loop post-market surveillance (U.S. Food and Drug Administration, 2025; European Union, 2024; National Medical Products Administration, 2026). All three expect a predetermined, pre-cleared change-control mechanism rather than a fresh review for every model update. The FDA framework is furthest along but still in draft: the AI-specific lifecycle guidance was issued for comment in January 2025 and has not yet been finalized (U.S. Food and Drug Administration, 2025); the EU equivalent is under active development at MDCG and IMDRF level (European Commission/MDCG, 2025); the NMPA risk-re-evaluation-on-change requirement takes effect March 2027. Convergence of this kind is rare, and it means an assurance architecture built once can serve every jurisdiction a manufacturer sells into.

The assurance side is already moving on the provider side, in the United States first. A leading academic medical center reviews its own clinical algorithms before internal use, and a large industry coalition of major US health systems and technology vendors has built a shared model-card registry (see references). A more ambitious companion effort, a nationwide network of independent assurance labs to test algorithms before deployment, was floated by that same coalition but was not delivered at scale; it was quietly abandoned in 2025, and the coalition's own founder now calls the concept a mistake (see references). Two gaps remain open. The model card is point-in-time, while all three regulators are converging on continuous lifecycle re-certification. And this progress is American and provider-side, while the 2027-2028 enforcement wave is European and Chinese, and the harder problem, an independent, pre-deployment testing gate, is the one that did not survive contact with reality. That is the window: the manufacturer-side, cross-jurisdictional gate is unbuilt, and the nearest attempt at it has already shown how hard it is.

4. Two futures for 2030

The methodology follows the Oxford Scenarios Programme (Ramírez et al., 2020). The source material is the UK Government Office for Science's AI 2030 scenario set, which maps five scenarios across five axes of uncertainty (Government Office for Science, 2025); this paper compresses that set to the two poles most decisive for a clinical-AI

portfolio, trust and adoption, and adapts them to a western regulatory environment (Lang & Ramirez, 2021). That compression is a deliberate simplification for tractability, not a claim that the source treats trust as the only axis that matters. Both scenarios describe the external environment only; removing any organisation's name leaves a plausible future intact, and both are treated as equally plausible for the purpose of stress-testing a strategy against each.

Frozen Trust. This is the world where clinical confidence in AI breaks before the technology matures. Trust goes first. The rest follows. By 2030 a run of high-profile AI-assisted clinical errors has triggered a regulatory reaction: moratoriums on new AI-enabled device approvals across the EU and US, health systems retreating from integrated deployments to isolated single-tool pilots, and mandatory human sign-off on every output. The sign-off restores caution at the price of the efficiency gains early adopters had banked. Most clinicians are under-prepared to supervise the tools they are handed; the training never scaled with the technology.

AI keeps improving. That stops mattering. Capability is no longer the constraint; trust and regulation are, and the binding test becomes proof of safety rather than speed of deployment: can the manufacturer demonstrate a pre-cleared, auditable change pathway, or does every update reopen the file. Power moves two ways at once, toward regulators and the institutions that can demonstrate assurance, and offshore, toward vendors in less-regulated markets who scale on larger clinical datasets and enter western markets through partnerships that sidestep direct regulatory exposure. The scarce asset is no longer the best model. It is an auditable one.

Full Commitment. This is the opposite bet, and it pays off. By 2030 the clinical evidence is in: AI-enabled tools have measurably cut misdiagnosis rates and improved outcomes in cardiology, radiology, and intensive care. The argument about whether AI works is over. The argument about who controls it is not.

Integration becomes the battleground, not model performance. The health systems and vendors that built integrated ecosystems early, and paid for the people to run them, lock in platform relationships and outcome-data advantages that latecomers cannot buy back. Regulation adapts instead of resisting, shifting from pre-market certainty to post-market surveillance, so the binding test becomes the speed of safe deployment. Power concentrates wherever the integration layer and the outcome data sit, which means with the platforms that stand between the model and the clinician. Value follows them, migrating from the model itself to the governed pipeline around it.

5. The plays

Four questions structure the response to each scenario: what AI changes about why a medtech company exists in clinical care (Purposes), who gains power (Players), which alliances matter and which dependencies turn risky (Partnerships), and what must change once the model ships (Processes). Within each scenario the four reinforce one another, a compounding effect set out in full after the four plays below.

Purposes. AI changes what a medtech company sells. Value moves from hardware and point algorithms to trusted, governed clinical decisions, and growth shifts from units shipped to recurring revenue from outcome-proven platforms. Under Frozen Trust the proposition that survives is safety: reposition the portfolio as AI under supervision, built on a Predetermined Change Control Plan for every product, cleared once at conformity assessment rather than re-litigated at every update, shipping with an explainability view, an audit log, and a documented override path signed off by a named Chief Medical Safety owner. One unexplained error resets the adoption curve; a pre-cleared change pathway lets a manufacturer keep iterating without resetting it again at the next release. Under Full Commitment, value shifts to the outcome-proven platform: drop the device-maker framing and price on outcomes, tying each

product P&L to misdiagnosis reduction and time-to-decision, with a scale-or-stop gate that pulls any model failing its post-market outcome threshold. Whoever prices on outcomes first captures a data flywheel the others cannot rebuild.

Players. The model matters less than whoever can certify it, integrate it, or refuse it. Under Frozen Trust power swings to regulators, hospital risk officers, and medical-liability insurers; they become the real buyers, not the clinical end-user. The play is to stand up a regulatory-and-evidence function as a front-line commercial unit, with named owners who co-author validation studies with clinical opinion leaders and sit in deals from day one. When approval and liability gate every sale, the team that can produce assurance is the team that closes business. Under Full Commitment power concentrates with the clinicians and health systems that set the integration standard, and with whoever holds the outcome data. The play is to convert lead clinicians into a governed champion network that decides what scales, backed by a data-rights clause in every deployment. Lose the data rights and a company rents its own advantage back.

Partnerships. AI turns a product company into one node in a clinical-AI stack, and dependencies that were manageable become dangerous. Under Frozen Trust the credibility-conferring relationships decide who trades at all: regulators, notified bodies, and the academic medical centres that can validate at scale. The play is to co-write post-market surveillance and validation standards with regulators and opinion leaders now, before those standards are set by committee, and to dual-source model supply so no single offshore supplier can strand a product line at the worst moment. Under Full Commitment the decisive alliances are the anchor health systems already in most incumbent portfolios. The play is to restructure them from co-development agreements into revenue-share co-investments in a vendor-neutral clinical AI validation standard, making anchor hospitals financial stakeholders in the standard that governs their own AI procurement. Ownership makes the alignment structural. Two dependencies stay dangerous in this future: ceding the orchestration layer to a distribution partner who then owns the customer relationship, and accepting product liability for third-party algorithms. The incumbent must operate the validation process and leave product risk with the manufacturer.

Processes. AI breaks the old product process at both ends. Before shipping, formulation stays exclusive while communication turns transparent. After shipping, pilots succeed and production stalls on the missing execution layer, the one that makes AI governable, cost-predictable, and accountable. Under Frozen Trust the play is to pre-build a regulatory-grade evidence pipeline so approval is never the bottleneck, design human sign-off as smart triage rather than blanket review, and treat data lineage as proof rather than afterthought. Under Full Commitment the core process turns continuous: a repeatable governed install across health systems, a silent-failure audit before any output informs a clinical decision, and the gap between the rules and what actually gets executed reported as a board metric.

Multiplicity: the four plays compound, in a specific order. Under Full Commitment the chain runs Partnerships to Purposes to Players to Processes: owning the vendor-neutral integration layer is what lets a manufacturer reposition as a governed clinical platform; that credibility is what mobilises clinical leaders into the champion network that decides what scales; and that network is what funds the continuous, governed install that keeps producing the outcome evidence the platform depends on. The four compound into a lasting scaling advantage rather than simply adding up. Under Frozen Trust the chain runs the other way, Processes to Players to Partnerships to Purposes: the regulatory-grade evidence pipeline and pre-cleared change-control pathway function as a governance moat; that moat is what lets clinician oversight operate as a genuine differentiator instead of a compliance cost; that standing is what earns a seat shaping the regulatory agenda before committees set it without input; and that influence is what

lets the portfolio compete on trusted, governed AI rather than defend it after the fact. The four compound into a resilient trust position. None of the plays wins a scenario by itself.

6. The robust core: five moves that win either way

Read across the scenarios rather than down them. Five moves recur in both futures, as a regulatory moat under Frozen Trust and as a scaling engine under Full Commitment. They are structurally valuable in different ways under each future. These belong in any clinical-AI plan now, before the signal lands.

No-regret move	In Frozen Trust	In Full Commitment
Governed execution	Auditable change pathway keeps products in market without resetting trust after every update	The governed deployment pipeline latecomers cannot rebuild
Clinician oversight	The assurance signal regulators and risk officers need: proof that clinicians remain in charge of the machine	Adoption lever that builds the clinical champion network driving scale
Accountability architecture	Answers who owns consequences before the first error occurs	Scales with the platform; its absence creates catastrophic liability at integration scale
Silent-failure discipline	Engineering-grade checking at inference is the assurance baseline regulators demand	Catches drift before it erodes the outcome data the platform monetises
Clinical validation platform	Cross-jurisdictional credential: one architecture for FDA, EU AI Act, and NMPA	Growth-engine candidate: compliance infrastructure that transitions to recurring revenue

7. Signals, breakpoints, and agility

The playbook only works if the watch is live. The signals are cheap to watch and expensive to miss; the recommendation is a standing signal review at every executive meeting where AI is on the agenda.

Toward Frozen Trust (activate defensive plays)	Toward Full Commitment (activate offensive plays)
Rising AI-related clinical incidents or near-misses	Accumulating peer-reviewed outcome proof
Regulatory or moratorium signals in any jurisdiction	Regulators shifting to post-market surveillance frameworks
Procurement pauses from hospital networks	Integration deals closing between AI vendors and hospital systems
Falling clinician trust scores or adoption reversals	Adoption rising across comparable product lines

Three breakpoints would invalidate the map itself. If a cluster of AI-assisted errors goes public or a liability ruling leaves clinicians personally exposed, trust snaps at once; offensive plays freeze and the governance moat becomes the product line. If mutual recognition across jurisdictions or a deregulatory turn cuts the cost of approval towards zero, the compliance moat thins and advantage shifts to deployment speed and distribution. If an open, publicly certified model stack makes validated AI a commodity, pricing power leaves the certificate and moves to outcome evidence and workflow lock-in. Signals say which future is arriving. Breakpoints say whether the map still holds.

Three agility levers decide whether an organisation can actually run this playbook. Strategic sensitivity: does leadership surface AI possibilities and translate them into action, at executive level, on a standing cadence?

Leadership unity: does stated priority translate into resource decisions below the top two layers, not just communications? Resource fluidity: is AI transition capacity ring-fenced, or do employees carry existing KPIs and the AI build simultaneously? In most incumbents the third lever is the gap, and it is a structural problem. The fix is to ring-fence transition capacity and rebalance KPIs to reflect the dual mandate.

8. Strategic synthesis

Two scenarios, four plays, one robust core, one constant purpose. Read together, Sections 5 through 7 compound into a single instruction: commit to the core now, and watch the signals.

Multiplicity. The four plays reinforce one another within each scenario rather than adding up independently: compounding into a lasting scaling advantage under Full Commitment, and into a resilient trust position under Frozen Trust, in the order set out in Section 5.

Common plays. Governed, accountable, human-led execution holds in both futures: clinician oversight differentiates, a named owner answers for every deployment before it ships, and a silent-failure audit runs before any output informs a decision. These are the no-regret moves in Section 6, and the point is to move them into the plan now, not after the signal lands.

Signals. The signals are cheap to watch and expensive to miss: rising incidents, moratorium talk, and procurement pauses toward Frozen Trust; accumulating outcome proof and closing integration deals toward Full Commitment. The full signal set and the three breakpoints that would invalidate the map itself are in Section 7.

Agility. The no-regret core runs in both futures while the scenario-specific plays sit pre-built, triggered by the same signals. An organisation that has ring-fenced transition capacity and aligned leadership below the top two layers can flip from defensive to offensive without re-planning from scratch.

Same doctor, same scan. The plays decide whether the experience from ten thousand cases is in the room with her.

9. The proposal worth sizing first

The fifth no-regret move deserves its own decision. Built on a vendor-neutral imaging or informatics layer that most incumbents already own, a clinical validation platform can operate as a conformity-assessment engine for clinical AI (Exhibit 1): algorithm vendors submit validation evidence through a guided, automated process; deployment sites provide post-market surveillance proofs; the operator periodically certifies algorithms against one shared standard aligned to FDA, EU AI Act, and NMPA requirements. The model already governs safety-critical industries worldwide: UL in the United States, TÜV and BSI in Europe, the China Compulsory Certification system. Applied to clinical AI, it converts compliance infrastructure that every manufacturer must build anyway into a certification credential that travels. Each new jurisdiction becomes localisation work on the same architecture. The liability architecture is the design constraint: the operator certifies the process, never the product, so product risk stays with the manufacturer and the deploying clinician. Cloud partners extend distribution reach without owning the clinical standard.

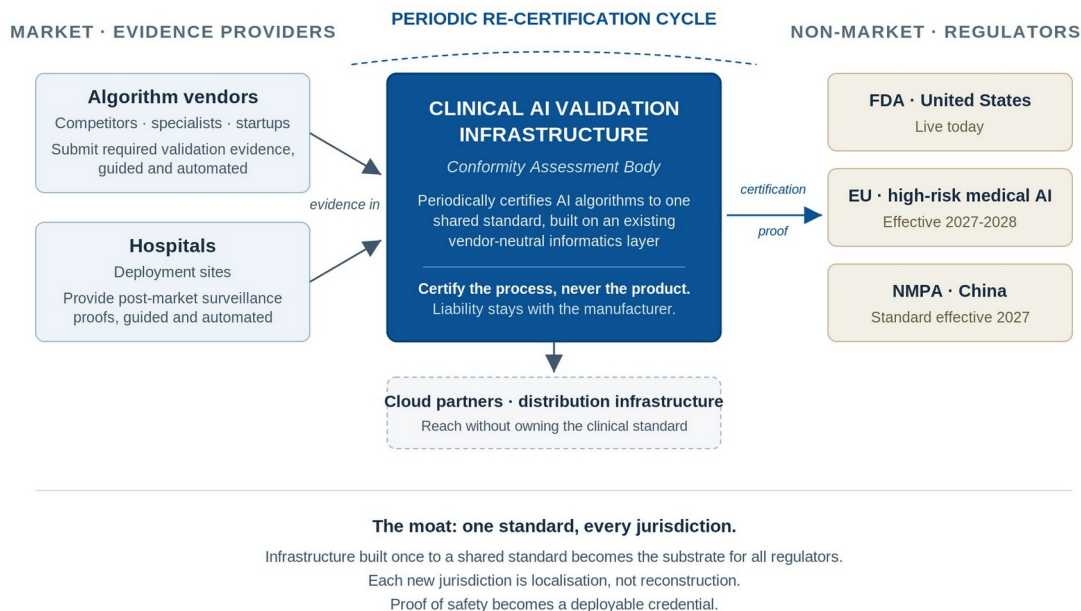


Exhibit 1: The validation platform as a conformity-assessment engine. Evidence flows in from vendors and deployment sites; certification proof flows out to FDA, EU, and NMPA; distribution partners extend reach without owning the standard.

The recommended first step is deliberately small: a 90-day workstream to size the validation platform as a business, covering addressable certification demand under the 2027-2028 EU and NMPA enforcement timelines, pricing against notified-body benchmarks, and the contractual architecture that keeps liability with the algorithm vendor. The question the workstream must answer is directional: whether clinical AI validation infrastructure can transition from a governance requirement into a recurring-revenue growth engine. If yes, the first mover defines the category. If no, the same infrastructure still protects the existing portfolio in either future, and nothing is lost but the study.

10. About this paper and its author

This paper is an open contribution. Any company that runs the playbook keeps full use of it; the scenario-to-plays methodology is replicable for any regulated-AI portfolio, in medtech and beyond.

I spent 25 years in global health technology, most recently leading AI transformation work at one of the sector's largest manufacturers, and developed this playbook on the Oxford AI-Driven Business Transformation Executive Programme at Saïd Business School. This edition was completed after leaving that role, and contains public data only. I am now looking for the next senior role where one person is accountable for an outcome like this, not a methodology. Selected writing and a full track record: bogdanrynkowski.com · LinkedIn: [rynkowski](#).

Appendix A: Frequently anticipated questions

Q: Why a scenario playbook instead of a forecast? A forecast bets on one future and fails quietly when the future moves. The playbook bets on neither scenario: the no-regret core runs now, and the scenario-specific plays sit pre-built, triggered by the signals in Section 7. The organisation flips from defensive to offensive without re-planning from scratch.

Q: Why these two scenarios, when the source set has five? The UK Government Office for Science AI 2030 set maps five scenarios across five axes, built for policymakers and publicly available (Government Office for Science, 2025). This paper compresses that set to the two poles built around the single axis that dominates every capital-allocation decision in a clinical-AI portfolio: whether clinical and regulatory trust in AI holds or fractures. That is a deliberate simplification, not a claim that the original five-scenario, five-axis set collapses to one dimension elsewhere. They remain deliberately external futures; no company appears in either, so they cannot flatter anyone's strategy.

Q: Doesn't certifying competitors' algorithms help competitors? It helps them clear a bar the operator sets. UL certifying a rival's appliance never weakened UL; it made UL the gate. Competitors submitting evidence to the operator's standard means the standard, the audit data, and the recurring certification revenue sit with the operator. The alternative is that a notified body, a cloud platform, or a provider-led assurance coalition builds the gate instead.

Q: Isn't this already being built by major players in the US? A US example shows both the promise and the limit of this pattern, and for a European reader it is a lesson worth learning at second hand rather than first. A leading US academic medical center runs real-world clinical validation research on its own multi-institutional data platform, and a large US industry coalition, backed by more than 3,000 member organisations including several major American health systems, built a real and growing AI model-card registry (see references). But that same coalition's headline proposal, a nationwide network of independent assurance labs to test algorithms before deployment, was quietly abandoned in 2025; its own founder now calls the idea a mistake, and the coalition pivoted to post-deployment governance tools instead (see references). That structural difference matters for reading across to this proposal: the coalition's effort was a multi-stakeholder, nonprofit-coordinated attempt to get thousands of members to converge on a shared standard, while the proposal here is a single-operator commercial model with one accountable owner. The two are not a like-for-like comparison. A single operator avoids the collective-action problem that likely contributed to the assurance labs stalling, but takes on a different risk in its place: without the neutrality a nonprofit, multi-stakeholder body can claim, a commercial operator has to earn independence credibility deal by deal, which is exactly what the liability wall in the next question is designed to protect.

For a European reader, the lesson is not that the US is ahead. It is that the most resourced, most motivated attempt yet at exactly this kind of cross-institutional validation infrastructure still could not get the pre-deployment testing model to work at scale, even with thousands of members and dedicated funding. Provider-led efforts of this kind are not a US-only pattern either: the UK's NHS AI Lab has run a comparable national programme for safe AI adoption since 2019, though its remit is policy and capability-building rather than algorithm certification (see references). None of these efforts closes the loop this proposal targets: a manufacturer-side, cross-jurisdictional gate aligned to FDA, EU AI Act, and NMPA requirements at once, which is precisely the gate a European manufacturer sitting inside the EU AI Act's own 2027-2028 enforcement window is well placed to build first. The abandoned assurance-lab concept is, if anything, the strongest evidence yet that the window is still open, and that the model which wins is likely to look different from the one that just failed.

Q: What about liability? The design constraint is fixed from day one: the operator certifies the process only. Product liability stays with the algorithm manufacturer and the deploying clinician. This is the conformity-assessment-body model that already governs safety-critical industries, and it is the reason the platform must never be positioned as a distributor or reseller of third-party algorithms.

Q: Can an incumbent afford this? The platform is built on assets most incumbents already own (a vendor-neutral imaging or informatics layer) and on validation work every portfolio must fund anyway for its own regulatory survival.

The increment is productisation of infrastructure that already exists. That is why the recommendation stops at a 90-day sizing study before any build commitment.

Q: What would falsify the playbook? Three breakpoints, named in Section 7: a trust snap that freezes all offensive plays, mutual recognition that cuts approval costs towards zero and thins the compliance moat, or an open certified model stack that commoditises validated AI. Each is watched alongside the scenario signals. A strategy that cannot name its own falsifiers is a bet.

Q: Why publish this openly? Because the methodology proves itself in the reading, and because the constraint it addresses is industry-wide. Any company that runs the playbook keeps full use of it. The author's background and current situation are stated plainly in Section 10.

Q: What is the single next step? Commission the 90-day sizing study on the validation platform. Everything else in the no-regret core is actionable within existing mandates and budgets.

Appendix B: References

- European Commission/MDCG (2025). MDCG 2025-6: Guidance on the interaction between the AI Act and the MDR/IVDR for medical device AI. Available at: <https://health.ec.europa.eu>
- European Union (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Available at: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- Government Office for Science (2025). AI scenarios 2030: helping policymakers plan for the future of AI. Available at: <https://www.gov.uk/government/publications/ai-scenarios-2030-helping-policymakers-plan-for-the-future-of-ai>
- Lang, T. & Ramirez, R. (2021). Getting the most from publicly available scenarios: 5 ways to avoid costly mistakes. California Management Review. Available at: <https://cmr.berkeley.edu/2021/04/getting-the-most-from-publicly-available-scenarios/>
- Mayo Clinic Platform (2024). How do you construct a safe, effective algorithm? Available at: <https://www.mayoclinicplatform.org/2024/11/13/how-do-you-construct-a-safe-effective-algorithm/>
- Mayo Clinic Platform (2025). Finding the best AI algorithms: is FDA approval enough? Available at: <https://www.mayoclinicplatform.org/2025/04/17/finding-the-best-ai-algorithms-is-fda-approval-enough/>
- National Medical Products Administration (2026). YY/T 1406-2026: Guidance on the Application of GB/T 42062 to Medical Device Software. Available at: <https://chinameddevice.com/nmpa-software-standard>
- Future Health Index (2026). AI in practice: the 2026 Future Health Index report. Available at: <https://www.philips.com/a-w/about/news/future-health-index/reports/2026/ai-in-practice.html>
- Partovi, S. (2026). AI is working in healthcare. Now let's scale it. Philips Innovation Matters (blog), 11 June 2026. Available at: <https://www.philips.com/a-w/about/news/archive/blogs/innovation-matters/2026/ai-is-working-in-healthcare-now-lets-scale-it.html>
- Ramírez, R., Lang, T., Selin, C. & Khong, C. (2020). Oxford Scenarios Programme. Oxford: Saïd Business School, University of Oxford.
- Rynkowski, B. (2026). Selected articles. LinkedIn. Available at: <https://www.linkedin.com/in/rynkowski/recent-activity/articles/>
- U.S. Food and Drug Administration (2025). Artificial Intelligence-Enabled Device Software Functions: Lifecycle Management and Marketing Submission Recommendations. Draft Guidance. Available at: <https://www.fda.gov/regulatory-information/search-fda-guidance-documents/artificial-intelligence-enabled-device-software-functions-lifecycle-management-and-marketing>
- Mayo Clinic Platform (2026). Accelerating AI innovation in healthcare: real-world clinical research applications on the Mayo Clinic Platform. npj Health Systems. Available at: <https://www.nature.com/articles/s44401-026-00068-1>
- Coalition for Health AI (2026). Model card registry. Available at: <https://www.chai.org/>
- Beavins, E. (2026). CHAI's AI oversight ambitions falter with scrapped AI labs. Fierce Healthcare. Available at: <https://www.fiercehealthcare.com/ai-and-machine-learning/inside-chais-failed-assurance-labs>
- Timmis, J. et al. (2025). A mixed methods formative evaluation of the United Kingdom National Health Service Artificial Intelligence Lab. Available at: <https://pmc.ncbi.nlm.nih.gov/articles/PMC12271305/>