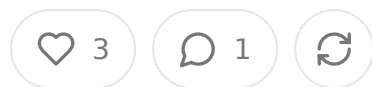


What actually matters in clinical AI right now: a reality check for health tech investors

NOV 28, 2025



Share

DISCLAIMER: The views and opinions expressed in this essay are solely my own and do not reflect the views, opinions, or positions of my employer, Datavant, or any affiliates.

If you are interested in joining my generalist healthcare angel syndicate, reach out to trey@onhealthcare.tech or send me a DM. Accredited investors only.

Abstract

This essay examines six recent publications that represent the most significant advances in clinical AI over the past six months. Rather than focusing on incremental improvements in narrow tasks or benchmarks, these papers demonstrate AI systems functioning in real clinical workflows with measurable impact on physician time, diagnostic accuracy, medication safety, and patient outcomes. The analysis reveals three critical investment themes: the shift from demonstration studies to pragmatic randomized trials, the superiority of human-AI collaboration over full automation, and the emergence of sophisticated evaluation frameworks that will shape regulatory pathways. For angel investors evaluating health tech opportunities, understanding these methodological shifts is as important as understanding the underlying technology, because they define which companies will successfully navigate clinical validation and health system adoption.

Introduction

The healthcare AI landscape has reached an inflection point where methodology matters more than model architecture. Over the past 18 months, we've seen a fundamental shift from papers showing that AI can do impressive things in controlled settings to papers demonstrating that AI can improve actual clinical outcomes in routine practice. This distinction is critical for investors because it separates technologies that will generate real revenue from technologies that will remain trapped in pilot purgatory.

The papers covered in this essay represent the current state of the art not because they achieve the highest benchmark scores or use the most sophisticated algorithms but because they answer questions that health system decision-makers actually care about. Can this technology reduce physician burnout in real-world conditions? Can it catch medication errors that human pharmacists miss? Does it introduce system bias that will expose us to liability? Can we deploy it without regulatory friction? These are the questions that determine whether a million-dollar pilot contract becomes a ten-million-dollar enterprise deployment, and these papers provide the evidence base that makes those decisions possible.

For angel investors, especially those without deep clinical backgrounds, understanding the methodological evolution of clinical AI research is essential because it tells you what kinds of companies will succeed over the next five years: companies that are building products aligned with these emerging evidence standards will navigate procurement processes smoothly. The companies that are still pitching based on impressive demos and favorable case studies will struggle.

The ambient AI scribe trial from UCLA

The UCLA ambient AI scribe study published in NEJM AI represents the first large-scale pragmatic randomized controlled trial of ambient documentation technology embedded in routine clinical care. Previous evaluations of AI scribes have mostly been observational studies where a hospital implements the technology and then measures changes in documentation time or physician satisfaction. Those studies consist

show positive results, which is unsurprising given that any new technology deployed with institutional support and physician opt-in is going to look good in the short run due to selection bias and Hawthorne effects.

What makes the UCLA study different is the rigor of the experimental design. They randomized 238 physicians across 14 specialties into three arms: usual documentation without AI assistance, Microsoft DAX, or Nabla. The randomization happened at the physician level rather than the encounter level, which matters because it controls for physician-specific factors like baseline documentation speed, typing proficiency, and personal workflow preferences. They tracked approximately 72,000 encounters over the study period and measured documentation time using actual EHR log data extracted from system timestamps rather than self-reported estimates from physician surveys.

That methodological choice alone is worth emphasizing because clinician self-reports about time usage are notoriously unreliable. When you ask physicians how long they spend on documentation, you get numbers that are colored by how burned out they feel, how much they dislike the EHR, and whether they're trying to make a point to administrators about workload. EHR log data doesn't have those biases. It just tells you when the note was opened and when it was signed, adjusted for time spent on other activities within the EHR during that window.

The headline finding was that Nabla users reduced documentation time by approximately 41 seconds per note compared to 18 seconds in the control group. After adjusting for baseline differences and potential confounders, this translated to a 37 percent larger reduction in documentation time for Nabla versus control, and that difference reached statistical significance. Microsoft DAX showed similar direct effects but didn't quite hit significance thresholds, possibly because of sample size because the effect size was genuinely smaller. Both AI tools were associated with modest improvements in validated burnout and cognitive workload measures, around 7 percent versus control, using the Stanford Professional Fulfillment Index and the Task Load Index rather than homegrown survey instruments.

The safety profile is equally important to understand. Clinically relevant inaccuracies in AI-generated documentation were described as occasional, and only one mild patient safety event occurred during the entire study that was potentially attributable to AI-generated content. Patient acceptance was high, with fewer than 10 percent of patients declining to have the AI scribe used during their visit. This is important because one of the concerns with ambient AI is that patients might find it creepy or intrusive to have their conversations recorded and processed by an algorithm, but empirical evidence suggests that concern is mostly theoretical.

For investors, the key signal here is not the exact magnitude of time savings, which is modest in absolute terms, but the fact that we now have RCT-level evidence demonstrating measurable workflow impact in routine conditions across multiple specialties. This fundamentally changes the value proposition from theoretical efficiency gains supported by vendor-provided case studies to demonstrated outcomes with quantified effect sizes published in a peer-reviewed journal. That distinction matters enormously in enterprise sales cycles because it shifts the burden of proof. Previously, health systems considering ambient AI deployment had to make decisions based on limited evidence and vendor claims. Now they have a template for what evidence looks like, and they're going to start demanding similar evidence from vendors.

The study also establishes important precedent for how these technologies should be evaluated going forward. Pragmatic randomized trials embedded in routine clinical operations are feasible and they provide much more credible evidence than retrospective observational studies. Vendors that can point to similar trial evidence or that are actively partnering with academic medical centers to generate it will have a significant advantage in competitive procurements. Vendors that continue to rely on pilot testimonials and before-after analyses are going to find themselves at a disadvantage as health system procurement teams become more sophisticated about evidence standards.

The other piece that's worth understanding is what this study doesn't tell us. It wasn't powered to detect rare safety events, so we don't know if there are low-frequency, high-severity errors that might emerge at scale. The follow-up period was relatively

short, so we don't know if the time savings persist over longer periods or if physicians eventually revert to baseline as they find workarounds or as the novelty wears off. The study was conducted at a single large academic medical center, so we don't know if the results generalize to community hospitals or small practices with different populations and workflows. These aren't criticisms of the study, they're just limitations that define where the evidence ends and where judgment begins.

From an investment perspective, the companies operating in this space that look attractive are the ones that have strong clinical partnerships, that are actively generating peer-reviewed evidence, and that have a clear path to demonstrating value in ways that align with how health systems make purchasing decisions. The ones that are still selling primarily on user enthusiasm and anecdotal feedback are going to struggle as the market matures.

It's also worth noting that 40 seconds per note may not sound like much, but when you multiply it across millions of encounters per year in a large health system, it translates to meaningful FTE savings or alternatively to more time available for patient care. The business case for ambient AI has always been strong in theory, now we have credible evidence that the theory holds up in practice, which means adoption curves are likely to accelerate.

Medication safety and the case for human-AI collaboration

The medication safety paper published in *Cell Reports Medicine* is probably the methodologically important study from the past year on how to deploy large language models in high-stakes clinical environments, and its importance stems specifically from what it doesn't recommend. The study evaluated LLM-based clinical decision support for identifying drug-related problems across 16 specialties using 91 common error scenarios spanning 40 vignettes. The research team built and validated five RAG-enhanced LLM variants, selected the best-performing model, and then tested it in three distinct deployment modes: LLM-only CDSS, pharmacist with LLM copilot, and pharmacist alone.

The copilot mode, where pharmacists had access to the LLM's recommendations alongside their own clinical judgment, achieved 61 percent accuracy with an F1 around 0.59. More significantly, for drug-related problems with serious harm potential, the copilot mode was 1.5 times more accurate than pharmacists working alone. That's a substantial improvement in exactly the cases where accuracy matters most. The LLM-only mode, by comparison, underperformed the copilot configuration by enough that you wouldn't want to deploy it in production even if regulatory frameworks and liability considerations allowed it, which they mostly don't.

This finding aligns with essentially every other rigorous evaluation of clinical AI conducted over the past two years: human-AI collaboration consistently outperforms either humans or AI operating independently. The pattern holds across different clinical domains, different model architectures, and different task complexities. If you give clinicians access to high-quality AI recommendations and let them integrate those recommendations with their own expertise and contextual knowledge, you get better outcomes than you do from either component alone. This isn't a minor technical detail, it's a fundamental insight that should shape how every clinical AI company thinks about product design and deployment strategy.

For investors, the implication is that you should be extremely skeptical of any health tech company pitching full automation of clinical decision-making. Not because the underlying technology isn't impressive, it often is, but because the performance ceiling for collaboration is demonstrably higher than the performance ceiling for automation, and health systems understand this now. The companies that will win in clinical AI are the ones building for collaboration from day one, not the ones building autonomous systems and then grudgingly adding human review when customers or regulators push back.

The methodology here is also worth understanding in detail because it's likely to become a template that other researchers and companies follow. The research team used retrieval-augmented generation architecture rather than simply prompting the base model with clinical scenarios. That matters enormously because it allows the system to pull in current drug interaction databases, formulary information, and clinical guidelines without requiring constant retraining of the underlying model.

RAG architectures separate the knowledge base from the reasoning engine, which means you can keep the knowledge current without the computational expense and engineering overhead of continuous model retraining.

They also explicitly compared multiple deployment modes rather than just reporting aggregate accuracy numbers. That comparison is critical because it directly addresses the question that health system decision-makers actually care about: should we deploy this with human oversight or without it? The answer from the data is unambiguous: you should deploy it with human oversight, and specifically you should deploy it in a way that presents the AI's recommendations to trained clinicians who can evaluate them in context and override them when necessary.

The paper includes a detailed flow diagram showing how the copilot system could be integrated into existing pharmacy workflows. That kind of implementation-focused documentation signals that the research team was thinking about real-world deployment from the beginning rather than just trying to publish impressive benchmark results. It's also exactly the kind of documentation that health systems need when they're evaluating whether to implement a new technology, because it helps them understand what changes to workflows, staffing, and infrastructure are required.

For companies building in this space, the implication is that your product roadmap needs to assume human-in-the-loop deployment as the default architecture, not an afterthought. The user interface design, the workflow integration points, the training requirements, the feedback mechanisms for capturing clinician overrides and corrections, all of these need to be optimized for collaboration rather than automation. That's a harder product problem to solve than building a model that generates good recommendations in isolation, but it's the problem that actually matters when it comes to how health systems want to deploy this technology.

The other technical detail that's important is the use of RAG architecture. If you're evaluating a company that's building clinical decision support tools on top of a large language model without retrieval augmentation, that should raise questions about how they're handling hallucinations, outdated information, and knowledge gaps.

RAG isn't a magic solution to these problems, but it's increasingly becoming table stakes for production clinical AI systems because it provides a principled way to ground model outputs in curated knowledge sources.

The challenge with RAG is that you need high-quality, structured, continuously updated knowledge bases to retrieve from, and building or licensing those is expensive and time-consuming. For medication safety specifically, you need drug interaction databases, pharmacokinetic data, formulary restrictions, institutional protocols, and ideally patient-specific information about allergies, renal function, hepatic function, and concurrent medications. Integrating all of that into a RAG pipeline that runs fast enough to be useful in clinical workflows is non-trivial engineering work.

Companies that have solved those integration challenges or that have partnerships giving them access to high-quality structured medical knowledge are going to have a significant advantage. Companies that are trying to rely purely on the knowledge embedded in pretrained models are going to struggle with accuracy, safety, and reliability in production environments.

It's also worth noting that this study used curated vignettes rather than real patient data from live clinical workflows. Vignettes allow for careful control of the scenario and ensure that you're testing specific kinds of drug-related problems, but they don't fully capture the messiness of real clinical data where information is incomplete, contradictory, or buried in unstructured text. The performance numbers from this study should be viewed as an upper bound on what you'd expect in production deployment, not as a guaranteed baseline. Companies making claims based on vignette studies without real-world validation should be viewed with appropriate skepticism.

The broader message for investors is that medication safety is one of the highest-value applications for clinical AI because medication errors are common, expensive, and often preventable. The literature suggests that somewhere between 5 and 10 percent of hospitalized patients experience adverse drug events, and a substantial fraction of those are preventable. If you can build a system that reduces preventable adverse

events by even 20 or 30 percent, the value proposition is immediately obvious to hospital quality and safety teams. But you can only capture that value if you build the system correctly, which means human-AI collaboration, RAG architecture, robust knowledge bases, and continuous monitoring infrastructure.

Reasoning models in critical care diagnostics

The DeepSeek-R1 study published in Critical Care represents one of the first prospective evaluations demonstrating that reasoning-style large language models meaningfully improve diagnostic performance in high-acuity clinical settings. The study recruited 32 critical care residents from six tertiary hospitals and randomized them to diagnosing 48 diagnostically challenging critical illness cases either with or without access to the model's output. These were NEJM-style case challenges and similar complex cases published after the model's training cutoff date, which eliminates the possibility of training data leakage where the model has seen the cases during pretraining.

The model alone, operating without human input, achieved 60 percent top-1 diagnostic accuracy across the 48 cases. That's impressive given that these are specifically selected to be diagnostically challenging and that they span multiple organ systems and disease processes. For context, human critical care residents without AI assistance achieved 27 percent top-1 accuracy on the same cases. When residents had access to the AI's output, their accuracy jumped to 58 percent, nearly matching the model's solo performance. Diagnostic time roughly halved, dropping from a median of 1920 seconds to 972 seconds when residents used AI assistance.

Those are substantial improvements in both accuracy and efficiency, and important because they occurred with subspecialty residents who already have significant domain expertise, not with medical students or general practitioners who might be expected to benefit more from decision support. The magnitude of the effect suggests that reasoning models are providing genuine diagnostic insight rather than just serving as sophisticated search engines or differential diagnosis generators.

For investors, this study is significant for several reasons. First, it provides empirical evidence that reasoning models, the ones that explicitly generate step-by-step thought processes before producing final answers, are genuinely better at complex diagnostic tasks than standard instruction-tuned models that go directly from input to output. That matters because reasoning models are computationally more expensive to run. They require more tokens, more processing time, and therefore more infrastructure cost per query. You need evidence that the extra computational expense translates into meaningful performance gains that justify the cost, and this study provides that evidence at least for diagnostic reasoning in critical care.

Second, the magnitude of the improvement is large enough that it's difficult to dismiss as a marginal gain. Going from 27 percent to 58 percent diagnostic accuracy on challenging cases is the kind of delta that changes clinical outcomes, not just workflow efficiency. Diagnostic errors in critical care settings are high-stakes because patients are physiologically unstable and delays or errors in diagnosis can lead to rapid deterioration. If you can deploy a system like this in ICUs and it prevents a small number of diagnostic errors per year, the value proposition becomes immediately obvious to hospital administrators, quality improvement teams, and intensivists.

Third, the study followed TRIPOD-LLM reporting guidelines and conducted external consistency checks on model output, which signals where evaluation norms are heading for clinical AI research. TRIPOD-LLM is an extension of the TRIPOD framework for reporting prediction model studies, specifically adapted for LLM-based systems. It requires documentation of the prompt engineering process, reporting of consistency across multiple runs with identical inputs, separation of performance evaluation from prompt optimization, and transparency about model selection and hyperparameters. As more journals adopt these standards, companies that can demonstrate adherence will have advantages in publication, credibility, and regulatory pathways.

The limitation, which the authors explicitly acknowledge, is that this remains a vignette-based study rather than a live deployment with actual ICU patients in real time. Vignettes can't fully capture the messiness of real clinical data: the incompleteness of the history, the ambiguity of the findings, the pressure of the clock.

information that arrives in fragments over hours, the contradictory lab values that don't make physiological sense, the patients who can't communicate their symptoms clearly, the family members providing unreliable histories. All of those real-world complexities can degrade model performance relative to what you see in clean vignettes.

That said, the performance gap between residents with and without AI assistance is large enough that even if real-world results are 20 or 30 percent lower than vignette results, you'd still expect meaningful clinical benefit. That makes this a strong candidate for follow-on implementation trials, and I'd expect to see those studies starting within the next 12 to 18 months at academic medical centers that have the research infrastructure and institutional review board processes to support them.

From a technical perspective, the study demonstrates that reasoning models can be effective without fine-tuning on medical data, which is significant because fine-tuning is expensive and requires large labeled datasets that often don't exist for specialized clinical tasks. If you can get good performance from clever prompting and chain-of-thought reasoning with a base model, that substantially lowers the barrier to deploying these systems across different clinical domains.

For companies building diagnostic decision support tools, the implication is that reasoning models should be seriously considered for complex diagnostic tasks where the cognitive process involves synthesizing information from multiple sources, generating differential diagnoses, and weighing evidence for and against different possibilities. Those are exactly the kinds of tasks where explicit reasoning chains are likely to outperform black-box predictions. The tradeoff is computational cost, but as inference costs continue to decline, that tradeoff becomes more favorable over time.

The other important signal from this study is that diagnostic decision support is one of the highest-value applications for clinical AI because diagnostic errors are both common and consequential. Estimates suggest that around 10 to 15 percent of diagnoses in acute care settings are wrong or delayed, and those errors contribute to substantial morbidity, mortality, and cost. If you can build a system that reduces diagnostic error rates by even a modest percentage, the value creation is enormous.

But you have to get the deployment model right, which means human-in-the-loop, clear presentation of reasoning and uncertainty, and integration into clinical workflows that don't create alert fatigue or cognitive burden.

It's worth noting that the critical care setting is particularly well-suited for this of AI deployment because intensivists are already accustomed to working with decision support tools, they have time to review AI-generated recommendations because critical care workflows involve frequent reassessment, and the patient population is monitored continuously so errors can be caught relatively quickly. Those characteristics make ICUs a natural testbed for diagnostic AI, and successful deployments in ICU settings will likely pave the way for broader adoption in other acute care environments.

Multi-task clinical workflows and the future of triage systems

The npj Digital Medicine paper from Gaber and colleagues takes a different methodological approach than the other studies covered here, and that difference is instructive for understanding where clinical AI evaluation is heading. Instead of testing a single narrowly defined task like generating documentation or identifying drug interactions, they built Claude-based workflows that simultaneously predict Emergency Severity Index triage level, mapped patients to appropriate specialty referrals, and suggested likely diagnoses. They used 2,000 cases derived from MIM-III ICU data with structured vitals, history of present illness, and demographic information, then benchmarked several Claude-family models including a RAG-assisted Claude 3.5 Sonnet configuration.

The study is slightly older than the strict six-month window but it's too important to exclude because it represents a fundamental shift in how clinical AI systems are conceptualized and evaluated. Most clinical AI research focuses on single endpoints: can you predict sepsis, can you identify diabetic retinopathy on fundus photos, can you generate an accurate radiology report. That narrow focus makes sense for research purposes because it's easier to isolate what the model is doing well or poorly at.

when you're measuring performance on a single well-defined task. But it doesn't reflect how these systems will actually be used in clinical practice, where you do want three separate tools that each do one thing, you want a single integrated interface that helps with triage, referral, and preliminary diagnosis as part of a coherent clinical workflow.

For investors, this distinction matters because it changes how you think about product scope, competitive positioning, and defensibility. If the future of clinical AI is integrated multi-task systems rather than point solutions, then companies that have built narrow single-purpose tools are going to struggle unless they can expand their capabilities quickly and coherently. Conversely, companies that are already thinking about multi-task architectures and can demonstrate that their models handle multiple clinical reasoning steps in a coordinated way will have significant advantages in enterprise sales.

The study found that RAG-enhanced models outperformed base models across all three tasks, which aligns with the pattern we've seen in the medication safety review and elsewhere. Multi-task clinical reasoning requires access to current clinical guidelines, institutional protocols, specialist availability, local epidemiology, and other contextual information that can't be baked into a pretrained model. RAG provides a principled architectural approach for grounding model outputs in this kind of structured contextual knowledge without requiring continuous retraining.

The implication is that the barrier to entry for building competitive clinical AI products is rising because you need not just a good foundation model but also the infrastructure to integrate it with real-time data sources in ways that are fast, reliable, secure, and clinically validated. Small companies with limited engineering resources are going to struggle to build and maintain that infrastructure, especially as health systems start demanding evidence that the RAG pipelines are pulling from high-quality sources and that the retrieval mechanisms are robust to missing or contradictory data.

This paper also functions as a blueprint for product development in ways that are somewhat unusual for academic research. It includes a detailed workflow diagram

showing how the system processes inputs and generates outputs, specific implementation details about prompt structure and RAG configuration, and evaluation methodology that other teams could replicate. Any competent engineering team could use this paper as a starting point for building a similar system, which means the value in companies working on this problem is less about the core technology, which is increasingly commoditized, and more about execution, clinical partnerships, regulatory strategy, data access, and go-to-market.

The MIMIC-derived case methodology is both a strength and a limitation. MIMIC data is detailed, well-structured, and includes rich clinical information from ICU patients at a major academic medical center. That makes it excellent for testing whether models can perform complex reasoning tasks when given complete information. But ICU patients are sicker and more complex than typical emergency department patients, the documentation is more thorough, and the clinical context is quite different from what you'd see in a community hospital ED at 2 AM on a Saturday night. So the performance numbers from this study should be interpreted as evidence that the approach works in principle, not as a guarantee of what you'd see in production deployment across diverse clinical settings.

For companies targeting the emergency department triage and decision support market, this study provides a roadmap for what good looks like in terms of multi-task evaluation and RAG integration. It also highlights that triage systems need to do more than just assign severity scores, they need to route patients to appropriate specialties and provide preliminary diagnostic hypotheses that can guide initial workup. Single-task triage models that only predict ESI level are solving an incomplete problem. Real-world health systems are going to prefer integrated solutions that address the full workflow.

The other insight from this work is that evaluation frameworks need to move beyond single-task accuracy metrics to assess how well models perform across multiple related tasks simultaneously. If you're building a triage system that predicts severity, specialty, and diagnosis, you need to evaluate not just whether each individual prediction is accurate but whether the combination of predictions is clinically coherent. A model that predicts high severity, routes to cardiology, but suggests

dermatologic diagnosis is doing something wrong even if each individual prediction might be defensible in isolation.

That kind of coherence checking is difficult to formalize but it's critical for real world deployment. Companies that are thinking seriously about multi-task clinical need to build evaluation frameworks that can detect incoherent or internally contradictory outputs, and they need to have mitigation strategies for when those occur. That might mean adding consistency checks in the inference pipeline, using multi-step reasoning where later steps can catch errors in earlier steps, or flagging low-confidence cases for human review.

Sociodemographic bias and why fairer testing matters for enterprise adoption

The Nature Medicine bias paper from Mount Sinai is arguably the most widely cited clinical AI paper from the past year and its influence extends well beyond the academic research community into health system governance processes, vendor procurement requirements, and regulatory policy discussions. The study generated over 1.7 million AI recommendations by taking 1,000 emergency department vignettes and replicating each one with 32 distinct sociodemographic profiles across nine different large language models. They then systematically analyzed patterns in triage priority, diagnostic imaging recommendations, treatment aggressiveness, and mental health escalation decisions to identify whether recommendations differed based on patient demographics rather than clinical presentation.

What they found is that some models escalated care, especially mental health evaluation, primarily based on sociodemographic characteristics rather than clinical status. High-income profiles received recommendations for advanced imaging more frequently than low-income profiles despite identical symptom presentations. Low-income profiles were more likely to receive recommendations that no further testing was needed. These aren't small effects occurring at the margins of the distribution; they're systematic patterns that showed up consistently across multiple models and multiple clinical scenarios.

The study also found that demographic bias wasn't uniform across all clinical situations. Some scenarios showed minimal bias while others showed substantial disparities, and the magnitude and direction of bias varied across different models. That heterogeneity is important because it suggests that bias isn't simply a property of training data that affects all models equally, it's also influenced by model architecture, fine-tuning approaches, and prompt engineering. That means bias is potentially mitigable through careful model design and evaluation, but it also means you can't assume that a model is fair just because it was trained on diverse data.

For investors, this paper is critically important not because it suggests you should avoid companies working on clinical AI, that would be a profoundly misguided takeaway, but because it defines what due diligence questions you need to be asking and what red flags you need to watch for. If a company is building clinical decision support tools and they don't have a clear articulation of how they're testing for sociodemographic bias, that's a significant red flag. If they're relying exclusively on aggregate accuracy metrics without any subgroup analysis, that's a red flag. If they can't describe what their fairness testing framework looks like or what mitigation strategies they have in place for disparities they identify, that's a red flag.

The paper proposes an AI assurance framework that involves systematic stress testing of LLMs across sociodemographic axes before deployment. This is rapidly becoming a standard expectation in enterprise procurement processes, especially at large academic medical centers and integrated delivery networks that have sophisticated health equity programs and institutional commitments to reducing disparities. Companies that can demonstrate they've conducted this kind of stress testing and have documented mitigation strategies for any disparities they identified will have significant competitive advantages in sales cycles.

The liability and risk management implications are also substantial. If a health system deploys a clinical AI tool that systematically undertreats or inappropriately escalates care for certain demographic groups, they're exposed to discrimination claims under civil rights law, medical malpractice liability for patients who are harmed by inappropriate recommendations, regulatory scrutiny from CMS and state health departments, and reputational damage that can affect patient volumes and community

trust. Hospital risk management and legal teams are increasingly aware of these exposures, which means they're requiring vendors to provide detailed evidence of fairness testing as part of the contracting process.

Companies that can't provide that evidence aren't getting past the contracting stage regardless of how impressive their core technology is. This isn't theoretical, I'm seeing procurement processes at major health systems where fairness documentation is a mandatory requirement in the RFP, and vendors without it are eliminated before technical evaluation even begins. That's a fundamental shift in how clinical AI is being purchased, and it happened remarkably quickly once this paper and a few others like it demonstrated the scope of the bias problem.

One thing that sometimes gets lost in discussions about AI bias is that this isn't purely a social justice issue, although it absolutely is that, it's also a business risk issue and a clinical safety issue. Biased recommendations lead to worse patient outcomes for the affected populations, which leads to higher readmission rates, lower performance on quality metrics, lower patient satisfaction scores, and potentially lower reimbursement under value-based care contracts. Health systems care about equity because it's the right thing to do, but they also care about it because biased care is expensive care that produces bad outcomes and creates legal exposure.

The methodology of the study is worth understanding because it's likely to become a template for how fairness evaluation is done going forward. Creating synthetic vignettes with systematically varied demographic profiles allows you to isolate the effect of demographics on model recommendations while holding clinical presentation constant. That's difficult to do with real patient data because clinical presentation and sociodemographic characteristics are often correlated in ways that are hard to disentangle. A patient's ZIP code might be associated with different disease prevalence, different access to preventive care, different health literacy, and different clinical trajectories, not just with discriminatory treatment recommendations.

Synthetic vignettes let you ask the counterfactual question: if this exact patient with this exact clinical presentation had a different race or income level or insurance

status, would the model make different recommendations? That's a powerful tool for detecting bias, and it's something that companies should be doing routinely as part of their model validation process. The challenge is that creating high-quality synthetic vignettes requires clinical expertise to ensure they're realistic and that the demographic variations you're testing are clinically appropriate.

For companies building clinical AI products, the actionable implications are that you need to build fairness testing into your development and validation workflows from the beginning, not as an afterthought once someone raises concerns. That means maintaining demographic metadata in your training and validation datasets, conducting routine subgroup analyses, implementing bias detection algorithms, having clear escalation processes when disparities are identified. It also means being transparent with customers about what fairness testing you've done, what you found, and what mitigation strategies you implemented.

The companies that will be successful in clinical AI over the next five years are the ones that treat fairness as a core product requirement rather than a compliance checkbox. That requires investment in diverse training data, sophisticated evaluation frameworks, ongoing monitoring infrastructure, and partnerships with health equity researchers who can help identify blind spots and validate approaches. It's not clear and it's not easy, but it's increasingly non-negotiable for enterprise adoption.

Foundation models in radiology and the case for vertical specialization

The Radiology review from Tavakoli and colleagues is less of a primary research contribution and more of an agenda-setting synthesis document, but it's worth paying attention because Radiology is one of the most influential specialty journals and reviews published there tend to shape how academic radiology departments, imaging AI companies, and healthcare organizations think about technology adoption priorities. The paper synthesizes emerging evidence on foundation models and generative architectures in radiology with explicit focus on clinical readiness and workflow integration rather than just algorithmic performance metrics.

The central argument is that subspecialty-specific foundation models are likely to reach clinical readiness substantially faster than broad generalist models attempting to cover all of radiology or all of clinical medicine. That conclusion is grounded in several observations about radiology as a domain: it has relatively well-defined tasks with standardized imaging protocols, structured reporting requirements that create clean training data, quantifiable performance metrics that correlate with clinical utility, and established regulatory pathways for imaging AI through FDA's radiological device framework.

For investors, this review reinforces a thesis that's been gaining traction over the last 18 months: vertical specialization matters more in clinical AI than it does in consumer AI applications. In consumer contexts, there's a strong argument for building generalist models that can handle diverse and unpredictable user needs across many domains. In clinical applications, especially in procedural specialties like radiology, pathology, and dermatology, you're generally better off building models that are deeply optimized for narrow sets of tasks with rich domain-specific training data and tight integration into clinical workflows.

That doesn't mean there's no role for generalist clinical models. Ambient documentation tools are essentially generalist models and they're performing well because documentation is a relatively general task that doesn't require deep domain-specific knowledge. But when you're building tools that directly impact diagnosis and treatment decisions, depth matters more than breadth, and partnerships with subspecialists who can provide training data, validate outputs, and guide workflow integration become critical competitive advantages.

The companies that are winning radiology AI contracts at major health systems today are overwhelmingly the ones that have focused on specific imaging modalities and specific clinical indications: chest CT for lung nodule detection, mammography for breast cancer screening, brain MRI for stroke assessment. They have deep partnerships with radiologists in those subspecialties, subspecialty-specific training datasets that reflect current imaging protocols, integration with PACS and RIS systems that radiologists actually use, and clear value propositions around reducing turnaround time or improving diagnostic accuracy for specific high-value tasks.

Contrast that with companies trying to build general-purpose radiology AI that works across all modalities and all anatomical regions. Those companies face much harder training data challenges because they need representative data across everything, much harder validation challenges because performance requirements differ substantially across subspecialties, much harder go-to-market challenges because they're trying to be everything to everyone, and much harder competition because they're fighting specialist competitors in every niche simultaneously.

The review also emphasizes the critical importance of continuous monitoring and failure mode analysis for foundation models deployed in clinical settings. Foundation models are powerful but they're also opaque and they can fail in unexpected ways especially when they encounter input distributions that differ from their training data. A model trained primarily on outpatient imaging might perform poorly on patients where image quality is compromised by portable equipment, patient positioning constraints, and physiological instability. A model trained on academic medical center data might struggle with community hospital imaging where protocols and equipment vary.

Radiology departments that deploy these systems need robust infrastructure for tracking performance over time, flagging cases where model confidence is low or where recommendations differ substantially from radiologist interpretations, and feeding that information back into continuous improvement processes. That monitoring infrastructure is expensive to build and maintain, requiring dedicated MLOps engineering, clinical oversight, data pipelines, and analytics capabilities.

This creates another barrier to entry that favors larger, better-capitalized companies over smaller startups. You can't just train a model and ship it to customers, you need the entire operational infrastructure to support continuous monitoring, performance tracking, version control, regulatory compliance, and customer support. Small companies with limited engineering resources and runway are going to struggle to build and maintain that infrastructure at the quality level that enterprise health system customers require.

For imaging AI specifically, there are additional technical challenges around integration with existing radiology IT infrastructure. PACS systems, RIS systems, reporting workflows vary substantially across institutions and many are built on legacy technology stacks that make integration difficult. AI tools that require manual data export and import won't get adopted because radiologists won't tolerate workflow friction that slows down their reading. Tools that integrate seamlessly with existing workflows and present results within the radiologist's normal workspace achieve much higher adoption rates.

That means successful radiology AI companies need not just strong models but also strong integration engineering, partnerships with major PACS vendors, and deep understanding of radiology workflows and pain points. Those capabilities are hard to build and they take time to develop, which creates meaningful moats for companies that have them.

The review also touches on the role of generative AI in radiology beyond just discriminative tasks like lesion detection. Generative models can potentially help with report generation, protocol optimization, synthetic data generation for training, even image reconstruction from undersampled acquisitions. Those applications are still mostly in research phases but they represent substantial potential value if they can be validated and deployed safely.

For investors evaluating radiology AI companies, the key questions to ask are: What specific imaging modality and clinical indication are you focused on? Who are your radiology subspecialist partners and advisors? What training data do you have access to and how representative is it of real-world practice variation? How does your tool integrate into existing PACS and RIS workflows? What's your regulatory pathway and where are you in that process? What's your monitoring and performance tracking infrastructure? How do you handle model updates and version control? Those questions will quickly separate companies that have thought deeply about real-world deployment from companies that are still in the research demo phase.

What this means for health tech investors going forward

If you're evaluating health tech companies working on clinical AI in 2025 and beyond, the six papers covered in this essay point to several themes that should inform your diligence process and portfolio construction strategy. These aren't abstract academic observations, they're practical implications that directly affect which companies succeed in capturing enterprise value and which will remain trapped in pilot purgatory.

First, evidentiary standards are rising rapidly and methodology matters as much as technology. Companies that are partnering with academic medical centers to conduct pragmatic trials or that can point to peer-reviewed publications in respected journals are going to have substantially easier paths to health system adoption than companies still relying on pilot testimonials and vendor-provided case studies. The UCLA ambient scribe trial established a new benchmark for what good evidence looks like in this domain, and procurement teams at sophisticated health systems are going to be demanding similar quality evidence from vendors.

This creates a bifurcation in the market where well-funded companies with strong clinical research partnerships can generate the evidence needed to win enterprise contracts, while smaller companies without those resources struggle to differentiate themselves from competitors who have better evidence. For angel investors, this suggests you should be looking for companies that have strategic relationships with academic medical centers or integrated delivery networks that can serve as both development partners and validation sites. Companies that are purely product-focused without clinical research capabilities are going to find themselves at a disadvantage.

Second, human-AI collaboration is the deployment model that actually works in high-stakes clinical environments, not full automation. This has profound implications for product design, user experience, training requirements, workflow integration, and go-to-market strategy. If a company is pitching autonomous clinical decision-making, they're either naive about regulatory and liability constraints or they're deliberately

overselling what the technology can do and what customers want. Neither is a good sign.

The medication safety study and the critical care diagnostics study both demonstrate that collaborative systems where AI augments human judgment outperform either component working alone. That's not a temporary limitation that will disappear with better models, it's a fundamental insight about how complex cognitive tasks work in domains with high stakes and incomplete information. For investors, this means you should be evaluating whether companies have designed their products for collaboration from the beginning. Look at the user interface: does it present AI recommendations in ways that support human judgment rather than replacing it? Look at the workflow: are there clear mechanisms for clinician override and feedback? Look at the training materials: do they emphasize critical evaluation of AI output rather than blind trust?

Third, fairness testing and bias mitigation are rapidly moving from nice-to-have features to mandatory requirements for enterprise procurement. The Mount Sinai bias paper has had remarkable influence on how health systems think about clinical AI governance, and that influence is translating directly into RFP requirements and contract negotiations. Companies that can't demonstrate they've conducted systematic fairness testing across demographic subgroups are getting eliminated from competitive procurements before technical evaluation begins.

This is creating a new category of must-have capability that smaller companies lack. Fairness testing requires diverse training and validation datasets with demographic metadata, sophisticated evaluation frameworks, partnerships with health equity researchers, and ongoing monitoring infrastructure. That's expensive to build and maintain, which favors larger, better-capitalized companies. For angel investors, this suggests you should be asking detailed questions about fairness testing methodologies, what disparities have been identified, and what mitigation strategies are in place. Companies that can't answer those questions clearly are accumulating regulatory and liability risk.

Fourth, vertical specialization is increasingly important in clinical AI. The radio review makes this point explicitly, but it applies more broadly across clinical domains. Subspecialty-focused companies with deep clinical partnerships, narrow task focus, and rich domain-specific training data are generally going to outperform general approaches for diagnostic and treatment decision support applications. The exception is workflow tools like ambient documentation where the task is general enough that specialization doesn't confer major advantages.

For portfolio construction, this suggests a barbell strategy might make sense: invest in a few horizontal workflow tools that can deploy broadly across specialties, and invest in multiple vertical-specific diagnostic or decision support tools that go deep in particular clinical domains. Avoid companies in the middle that are trying to be somewhat specialized but not specialized enough to build real moats against focused competitors.

Fifth, the infrastructure requirements for deploying and maintaining production clinical AI systems are substantial and growing. Companies need MLOps capabilities, continuous monitoring infrastructure, integration engineering, regulatory expertise, clinical validation partnerships, fairness testing frameworks, and customer support operations. Small teams without significant engineering capacity and adequate resources are going to struggle to build all of that, which suggests the market will consolidate toward a smaller number of well-capitalized players who can afford to do things correctly.

This has implications for entry valuation and capital requirements. Clinical AI companies need more runway and more engineering headcount than equivalent consumer AI companies because the operational infrastructure requirements are demanding and the sales cycles are longer. Seed and Series A rounds that would be adequate for consumer AI products may be insufficient for clinical AI products, so investors need to underwrite accordingly.

Sixth, retrieval-augmented generation is becoming table stakes for clinical decision support applications. Multiple studies covered here demonstrate that RAG architectures outperform base models for tasks that require current medical

knowledge, institutional protocols, or patient-specific information. Companies that are trying to build clinical AI products on top of foundation models without retrieval augmentation are going to struggle with hallucinations, knowledge currency, and grounding.

But RAG introduces its own complexities. You need high-quality structured knowledge bases to retrieve from, you need fast and reliable retrieval mechanisms, you need security controls to prevent data leakage, and you need monitoring to catch when retrieval is pulling in incorrect or contradictory information. Companies that have solved those challenges or that have partnerships giving them access to structured medical knowledge sources have meaningful advantages. When you're evaluating companies, ask detailed questions about their RAG architecture: what knowledge sources are they retrieving from, how often are those sources updated, do they handle retrieval failures, how do they ensure retrieved information is relevant to the query.

Seventh and finally, regulatory pathways and reimbursement mechanisms are still evolving rapidly for clinical AI. FDA has established frameworks for certain categories of imaging AI and is developing frameworks for clinical decision support tools, but there's still substantial uncertainty about requirements for different types of systems. CMS has been cautious about creating specific reimbursement codes for AI-augmented services, which means most clinical AI is being paid for through institutional budgets rather than fee-for-service reimbursement.

Companies that have clear regulatory strategies and that are actively engaging with FDA and CMS are better positioned than companies that are ignoring regulatory questions until they become urgent. For investors, this means doing careful diligence on regulatory pathway and timeline, understanding what predicate devices or substantial equivalence claims exist, and having realistic expectations about how regulatory clearance and reimbursement establishment will take.

The broader pattern across all these papers is a shift from proof-of-concept research demonstrating that AI can perform clinical tasks to implementation science examining whether AI improves outcomes when deployed in real clinical

environments. That's a critical transition because it separates technologies that work in controlled conditions from technologies that work in the messy reality of actual healthcare delivery. The companies that understand this shift and are building products optimized for real-world deployment rather than benchmark performance are the ones that will capture the value being created in clinical AI over the next years.

For angel investors, especially those without deep healthcare backgrounds, the implication is that you need to do more operational and go-to-market diligence and less purely technical diligence than you might for consumer AI companies. The models are increasingly commoditized, the differentiation is in clinical partnerships, evidence generation, workflow integration, fairness testing, regulatory strategy, and operational infrastructure. Those are areas where domain expertise and execution matter more than algorithmic innovation.

The companies that will win in clinical AI are not necessarily the ones with the most impressive models or the highest benchmark scores. They're the ones that have figured out how to generate credible clinical evidence, build collaborative rather than autonomous systems, test systematically for bias, specialize deeply in valuable clinical domains, build robust operational infrastructure, and navigate complex regulatory and reimbursement pathways. Those capabilities take time, capital, and domain expertise to build, which creates meaningful barriers to entry and opportunities for well-positioned companies to establish durable competitive advantages.



3 Likes

← Previous

Next →

Discussion about this post

Comments

Restacks



Write a comment...



Stuart Miller 🌟 Haverin about... Nov 28

I like the summary Trey. It speaks volumes to what GOOD human product manage known for years. Understanding your market and the problem being sought to sol critical. Your proposed portfolio arrangement for VC's makes a lot of sense too wit balance across your 'barbell'. Sadly I see a number jumping in and placing multip on one type of technology.

Finally, I will repeat the assertion I have been making for some time.

Institutions and the IT and Finance leadership need to establish infrastructure and infrastructure management to organize, control, validate and apply performance and security checks against the AI tools landscape. The lack of creating this control and monitoring system will be the undoing of some healthcare organizations, across both payer and provider sectors. Management and Governance will become the controlling risk profile from the architectures and behaviors of any of these tools can be devastating and fatal.

♡ LIKE 💬 REPLY