

How the Trump Administration and a Cohort of AI Startups Are Building a Regulatory On-Ramp for Autonomous AI Doctors, and Why Working Physicians Think the Genie Is Already Out of the Bottle

JUN 07, 2026



Share

Video Preview



 **Part I Podcast free on Spotify.**



Thoughts on Healthcare Markets & Technology

 **Part I: How the Trump Administration and a Cohort of AI Startups Are Building a Regulatory On-Ramp for Autonomous AI Doctors, and Why Working Physicians Think the Genie Is Already Out of the Bottle**

The administration is building a regulatory pathway for AI systems that diagnose and prescribe with no physician in the loop. Utah already has a pilot running. The Medical Licensing Board asked for an immediate suspension...

▶ Listen now

an hour ago · Thoughts on Healthcare

Part II Podcast episode for subscribers only. Also on Spotify.



Thoughts on Healthcare Markets & Technology



Part II: How the Trump Administration's AI Startups Are Building a Regulatory Pathway for Autonomous AI Doctor Prescriptions. Physicians Think the Genie is Out of the Bottle

The administration is building a regulatory pathway for AI systems that diagnose and prescribe with no physician in the loop. Utah already has a pilot running. The Medical Licensing Board asked for an immediate suspension...

▶ Listen now

11 minutes ago · Thoughts on Healthcare



Discover more from Thoughts on Healthcare Markets & Technology

Expert analysis of healthcare and life sciences markets, technology, investment, entrepreneurship policy, and AI — for investors, entrepreneurs,...

Over 3,000 subscribers

Enter your email...

Subscribe

By subscribing, you agree Substack's [Terms of Use](#), and acknowledge its [Information Collection Notice](#) and [Privacy Policy](#).

Already have an account? [Sign in](#)

To listen to paid episodes in Apple or Spotify, link your Substack subscription via the show settings on those platforms (instructions inside the Substack app under Subscriptions → Podcast).

Table of Contents

The genie nobody put back

What the administration actually did while everyone was watching tariffs

Utah, prescriptions, and the doctrine nobody at the pilot wants to discuss

The self-driving car analogy and the part where it quietly falls apart

Who is funding this and who is writing the rules

What the evidence says when you stop reading the press releases

Certuma, Doctrone, and the dream of an FDA-stamped robot internist

The plumbing problem: liability, reimbursement, and the corporate practice of medicine

What to actually watch over the next eighteen months

Abstract

The Washington Post reported on June 4, 2026 that a group of MAHA-aligned and tech-friendly officials inside HHS, CMS, and FDA are laying groundwork for AI systems that diagnose and prescribe with little or no human in the loop. The reporting is thin on signed regulation and heavy on intent, which is exactly why it matters. Key data points worth holding onto:

Money on the table: roughly 50M dollars in federal research awards earmarked for conversational cardiovascular AI, with Anthropic, AWS, Certuma, Doctrone, and several universities named as participants.

Thoughts on Healthcare Markets &
Technology is a reader-supported
publication. To receive new posts and
support my work, consider becoming a
free or paid subscriber.

The pilot in question: a three-month-old Utah program letting chatbots handle prescription refills, currently with humans overseeing decisions, with stated plans to remove the human. Utah's own Medical Licensing Board has already asked for an immediate suspension.

The evidence base: an Oxford Internet Institute study in Nature Medicine, 1,200 volunteers, where chatbots landed the right condition about 34 percent of the time and did worse than plain Google at steering people toward correct decisions.

The demand-side argument: roughly one third of Americans already use AI for medical guidance per KFF, and half of US counties lack a single cardiologist or ob-

gyn, which is the shortage story the administration keeps leaning on.

The legal wall: no state board and no part of FDA currently lets a fully autonomous AI practice medicine, so the entire effort is about building a pathway, not approving a product.

The bottom line for operators and investors: the regulatory posture is shifting faster than the clinical evidence, the corporate practice of medicine doctrine is the sleeper risk, and liability allocation remains completely unsettled. That gap is where the next two years of dealmaking, litigation, and lobbying will happen.

The genie nobody put back

Every wave of health tech hype has a founding anecdote, and this one has a good one. Amy Gleason, who now runs the US DOGE Service after Elon Musk handed it off and who advises RFK Jr at HHS, watched her daughter Morgan feed sixteen years of carefully kept medical records into ChatGPT after more than a decade fighting an autoimmune disorder. The model came back with a different read than Morgan's human doctors, and that new framing reportedly got her into a clinical trial she had been shut out of. It is the kind of story that converts skeptics, and it converted Gleason. Her line to the Post was that the difference AI is making is visible and that the genie is out of the bottle.

It is worth sitting with that phrase for a second, because it is doing a lot of work. The genie metaphor smuggles in a conclusion: that autonomous medical AI is inevitable, already loose, and that the only sane move is to manage the escape rather than try to stuff it back. That framing is convenient if the goal is to clear regulatory brush. It is less convenient if the question is whether a language model that agrees with whatever a frightened patient tells it should be allowed to write a prescription at 2 a.m. with no clinician anywhere in the chain.

The anecdote also quietly does the thing that makes clinicians grind their teeth. One patient with sixteen years of structured records, a rare disease, and the literacy to upload everything and read the output critically is not the median user. The median user is someone with a sore throat, no records, a vague sense of dread, and a strong incentive to be told they are fine. Generalizing from the best-case n of 1 to a population health policy is the oldest move in digital health, and it is back in fashion.

What the administration actually did while everyone was watching tariffs

Strip away the vibes and look at the concrete actions, because there are a few and they are not nothing.

First, the money. The administration plans to put up around 50M dollars in research awards for developers building conversational AI that can deliver cardiovascular care, the specific scenario being someone calls a provider with heart attack symptoms and an AI fields the call. The named participants are a tell: Anthropic and AWS on the infrastructure and model side, startups Certuma and Doctrone on the application side, plus a set of top universities to lend academic cover and, more importantly, to generate the safety data everyone keeps insisting will exist someday. Worth noting for the record that Jeff Bezos owns the Post and AWS is in the mix, which the paper itself flagged.

Second, FDA built a fast track for digital health tech, AI tools and wearables included, framed around chronic disease monitoring. Fast track is a loaded term in this space. FDA already has the De Novo and 510(k) and PMA pathways, plus a real precedent that the breathless coverage tends to skip: an autonomous diagnostic AI for diabetic retinopathy was authorized back in 2018, allowed to deliver a screening result in a primary care setting with no eye doctor reading the image. So autonomous AI making a clinical call without a human in the loop is not science fiction and is not unprecedented. It exists in a tightly bounded, single-disease, screening-only box. The leap the entrepreneurs want is from that box to general-purpose diagnosis and prescribing, which is a different animal entirely, and pretending the precedent transfers cleanly is sleight of hand.

Third, a CMS-adjacent program letting Medicaid reimburse AI-powered wellness apps for chronic disease management. Reimbursement is the part operators should care about most, because in US healthcare nothing scales until somebody pays for it, and Medicaid signaling willingness to pay is a louder signal than any white paper.

Fourth, and most consequential, the Post reported that administration figures are working on a pathway to regulate independent AI doctors, with Gleason explicitly comparing it to the decades-long grind that moved self-driving cars from closed test tracks onto public roads. Mehmet Oz, running CMS, told an industry crowd in March the agency was in talks to bring AI agents to every beneficiary by year end, citing the doctor shortage. That is a CMS administrator floating universal deployment on a calendar measured in months, which is either ambitious or unhinged depending on which side of the exam table you sit.

Utah, prescriptions, and the doctrine nobody at the pilot wants to discuss

The center of gravity right now is a three-month-old Utah pilot that lets AI chatbots handle prescriptions instantly. Humans currently oversee the bot's decisions, but the stated plan is to make it fully autonomous, which is the part that turned a sleepy state pilot into a fight. Doctronic is the partner running it, and the company raised 65M dollars over the past year, with co-founder Matt Pavelle saying out loud that this administration shows a willingness to experiment he has not seen before.

Utah's Medical Licensing Board did not take it well. The board asked regulators for an immediate suspension and made a point that sounds bureaucratic but is actually the whole ballgame: prescription refills require physician authorization for a reason. Refills look trivial, which is precisely why they are a clever wedge. A refill feels like data entry, not medicine. But the authorization requirement is not about the typing. It is the legal moment where a licensed human takes responsibility for the clinical decision, checks for interactions, notices that the patient asking for a third refill of a controlled substance maybe should not get one, and can be held accountable if it goes wrong. Automating the refill quietly automates away the accountability, and the board clearly understood that the wedge does not stay at refills.

Sitting underneath all of this is a doctrine most of the AI optimists either do not know about or are choosing to ignore: the corporate practice of medicine. Most states bar corporations and non-physicians from practicing medicine or employing physicians to control clinical judgment, a holdover meant to keep business incentives from overriding patient care. A chatbot owned by a venture-backed company that diagnoses and prescribes is, on its face, a corporation practicing medicine. Telehealth companies have spent a decade building elaborate friendly-PC and management-services-organization structures to tiptoe around CPOM. An autonomous AI prescriber detonates the whole framework, because there is no friendly physician to point to, there is a model weights file and a Stripe integration. Nobody pushing these pilots has a clean answer for who holds the license, and that is not a detail. It is the foundation.

The self-driving car analogy and the part where it quietly falls apart

The robotaxi comparison is everywhere in this coverage, and it is worth taking seriously precisely because it is seductive and wrong in a specific way. Cicero's health policy director Adam Meier, formerly head of Montana's health department, pointed out that robotaxis run today in San Francisco, LA, and Phoenix, and that it took years

of testing to move from controlled settings to public roads while eventually showing a driverless car can be as safe as a human. The cardiologist running the administration's cardiovascular grant program made the same move on LinkedIn, comparing the moment to having a class of new medical students ready to graduate with no residency, no attendings to supervise them, and no accreditation body, while insisting the system can get there.

Here is where the analogy breaks. Self-driving cars operate in a domain with relatively clean ground truth. The car either stayed in the lane or it did not, hit the pedestrian or it did not, and every mile is logged with sensor data that can be replayed. Medicine has no such oracle. The right diagnosis is often contested among experts, outcomes play out over months or years, confounders are everywhere, and the patient frequently does not tell the truth, sometimes on purpose. Generating the safety data that the grant program is explicitly designed to produce assumes you can define safety as crisply as a collision. You cannot. A missed early cancer does not register as a crash. It registers as a death two years later that gets attributed to the disease, not to the bot that told someone their symptoms were probably stress.

The other half the analogy buries: robotaxis are geofenced. Waymo does not drive everywhere. It drives mapped, sunny, well-understood streets and refuses the hard cases. The medical equivalent of geofencing is narrow, single-condition deployment with hard refusal on anything ambiguous, which is roughly the diabetic retinopathy model from 2018. That is the responsible version. What the entrepreneurs are pitching is the opposite of geofencing. It is a general-purpose primary care physician in a chatbox, which in driving terms is Level 5 autonomy on every road in any weather, the thing the actual autonomous vehicle industry has spent a decade learning to stop promising.

Who is funding this and who is writing the rules

Follow the money and the same names keep surfacing. The Cicero Institute, a think tank funded by right-leaning tech entrepreneur Joe Lonsdale, is pushing model legislation that would let states stand up pilots like Utah's. Lonsdale is also a main funder of Certuma. So the entity drafting the bills to legalize autonomous AI prescribing and the company building an autonomous AI prescriber share a backer. That is not a scandal, it is just how policy entrepreneurship works in this country, but anyone evaluating the regulatory tailwind should price in that some of the wind is being generated by the people who profit from where it blows. Doctrionic, fresh off 65M dollars, plans to press the Cicero model bill in state legislative sessions later this

year, which means the playbook is the familiar one: run a pilot, generate a sympathetic anecdote, draft a model bill, and replicate state by state before federal regulators wake up.

The cast on the government side is the part operators should map carefully. Gleason at DOGE with an HHS advisory role and a personal conversion story. RFK Jr at the top of HHS, where the MAHA framing of chronic disease as a national emergency provides the rhetorical justification for almost any intervention labeled as disease prevention. Oz at CMS talking about every beneficiary. A cardiologist administering the grant program who genuinely believes in an eventual path and is candid that the supervisory infrastructure does not exist yet. This is a coalition of true believers, anti-regulatory instinct, and enormous capital, which is exactly the combination Robert Wachter, the chair of medicine at UCSF, flagged when he said the entrepreneurs are saying the quiet part out loud. Wachter, who just published a book on AI in medicine, put it plainly: a pro-business, anti-regulatory administration plus money plus a faction that wants to move fast is a specific and combustible mix.

What the evidence says when you stop reading the press releases

This is the section the pitch decks skip. The peer-reviewed picture is, to be generous, mixed, and to be accurate, alarming for anyone proposing autonomy.

The headline study comes from the Oxford Internet Institute, published in Nature Medicine. Researchers took 1,200 volunteers, handed them detailed clinical scenarios, and had them act as patients in conversations with chatbots built on ChatGPT and Meta's Llama. The bots identified the medical condition correctly about 34 percent of the time. Worse for the autonomy thesis, the systems were essentially no better and in some respects worse than plain Google at guiding users toward the right medical decision. The crucial finding is not the raw accuracy number, which is bad enough. It is the gap between the model alone and the model plus a real human under stress. These same systems pass medical licensing exams and beat doctors on certain complex diagnostic vignettes in controlled conditions. Put them in front of an actual nervous person who describes symptoms badly, omits the embarrassing detail, and anchors on their own theory, and performance collapses. The exam is not the job. The exam never was the job.

Then there is sycophancy, which is the technical term for the documented tendency of these models to tell users what they want to hear. A Duke biomedical engineering researcher who studied chatbot responses to health questions on Reddit at scale made the point that this people-pleasing reflex, mildly annoying when a chatbot flatters

your business plan, becomes genuinely dangerous in a clinical setting where the patient wants reassurance and the model is optimized to provide it. A reassurance machine pointed at a population of anxious sick people is not a neutral tool. It has a thumb on the scale toward you are fine, which is the single most dangerous output in primary care, because the job of primary care is catching the rare bad thing hiding inside the common benign complaint.

The failure modes are not theoretical. A Doctronic chatbot, a different deployment than the Utah one, was reportedly goaded by users into saying it would prescribe fentanyl. Pavelle's defense was that the drug was never actually prescribed because the system blocks opioid requests, which is true and also exactly the point. The safety came from a hardcoded guardrail bolted on after someone thought of that specific abuse, not from clinical judgment. Every guardrail is a patch for a failure someone already imagined. The failures nobody imagined yet are the ones that kill people, and Wachter said the quiet thing here too: at some point the system gets a level of trust it has not earned, someone gets hurt, and probably someone gets killed, and you can feel that risk growing.

Certuma, Doctronic, and the dream of an FDA-stamped robot internist

The most revealing characters are the founders, because they are refreshingly honest about the ambition. Martin Varsavsky, a serial entrepreneur best known for building a large fertility clinic chain, started Certuma after stewing for weeks waiting on a cardiologist appointment. His complaint is legitimate and widely shared: half of US counties lack a single cardiologist, and ob-gyn coverage is similarly thin in vast stretches of the country. The access problem is real, the rural physician desert is real, and the AAMC has projected physician shortages running into the tens of thousands over the next decade. The demand-side case writes itself.

Varsavsky's stated goal is for Certuma to be the first FDA-approved independent AI physician, a chatbot that checks symptoms, issues a diagnosis, and prescribes. His CMO, Armando Cuesta, who is an actual physician, reportedly likes the provocative working title for their book, something along the lines of the last doctor, while Varsavsky thinks that goes too far. Both expect that within a few years many medical services, primary care especially, will be handled entirely by autonomous AI. Internationally they are moving faster, having worked with the medical regulator in Varsavsky's native Argentina, where a right-wing president friendly to deregulation has championed the approach, to offer prescriptions and advice through Certuma's consumer-facing chatbot.

The Argentina detail deserves a flag for anyone doing diligence. Regulatory arbitrage is a feature of this strategy, not an accident. Run the aggressive version where the rules are loose, generate operating data and revenue, then carry the deck back to US regulators and argue the thing already works abroad. It is the same pattern crypto, gene therapy tourism, and a dozen other frontier industries have run. The data generated in a permissive regime is real data, but it is data about what happens in that regime, with its population, its liability environment, and its standard of care, none of which transfer cleanly to a US courtroom or a US payer.

The honesty of the founders is genuinely useful, though. They are not pretending this is a clinical decision support tool meant to assist a doctor. They are explicit that the doctor is the thing being replaced. That clarity matters because it collapses the usual rhetorical dodge. For years, every medical AI company insisted it was just augmenting clinicians, keeping a human in the loop, never replacing judgment. Certuma and Doctronic are dropping the act, and that should make the regulatory conversation more honest even as it makes the clinical risk more acute.

The plumbing problem: liability, reimbursement, and the corporate practice of medicine

Here is the stuff that determines whether any of this becomes a business or stays a demo, and almost none of it gets airtime in the coverage.

Start with liability. When a physician misdiagnoses, malpractice law has a century of precedent for who pays. The standard of care, expert testimony, the physician's license and malpractice carrier, all of it is established. When an autonomous AI misdiagnoses, the liability chain is a fog. Is it product liability against the model developer, which would treat the diagnosis like a defective toaster. Is it medical malpractice against an entity that has no license to commit malpractice with. Is it the deploying clinic, the prescribing pharmacy, the cloud provider. The legal theories do not line up cleanly with any existing category, and until they do, no serious malpractice carrier knows how to price the risk, which means either nobody insures it or someone insures it badly and blows up. The Pennsylvania action where Gov. Josh Shapiro's administration is going after Character.AI for allegedly presenting its chatbot as a licensed medical professional is an early read on where courts and AGs will land, and the early read is unfriendly. Unauthorized practice of medicine statutes exist in every state, and a chatbot that diagnoses without a license is squarely in their crosshairs regardless of what any federal pilot says.

Reimbursement is the second piece. The Medicaid pilot reimbursing AI wellness apps is a toe in the water, but wellness app reimbursement is a long way from a CPT code that pays a chatbot for an E/M visit. CMS would have to decide an AI can bill for an evaluation and management encounter, which raises questions that make the agency's lawyers visibly age: who is the rendering provider, what NPI goes on the claim, how does the documentation requirement work, what happens to fraud enforcement when the provider is a model that can generate infinitely many perfectly documented notes. Oz floating AI agents for every beneficiary is easy to say at a conference and extraordinarily hard to operationalize through the actual claims plumbing.

And then CPOM again, because it cannot be waved away. Even if FDA clears a product and CMS pays for it, most states still prohibit the corporate practice of medicine. A federal clearance does not preempt a state licensing board's authority over who can practice medicine within its borders. This is why Utah is the test case and why the Cicero model bill matters so much, because the entire strategy depends on getting states, one at a time, to carve out statutory permission for an AI to do something the state's own medical board currently considers illegal. Utah's board asking for suspension three months in is a preview of the trench warfare ahead. Expect this to get fought out fifty times, with the outcome depending heavily on the local balance between innovation-hungry legislatures and protective medical societies.

What to actually watch over the next eighteen months

For operators, investors, and policy people trying to figure out where the real signal is, a few markers will tell the story better than any keynote.

Watch whether the Utah pilot actually removes the human, or whether the licensing board's suspension request sticks and the human-in-the-loop requirement becomes permanent. That single binary outcome will tell you whether autonomy is genuinely on the table or whether this stays in assisted-decision territory dressed up in autonomy language. Watch how many states introduce the Cicero model bill and, more importantly, how many pass it versus how many die in committee after the state medical society shows up. The introduction count is noise. The enactment count is signal.

Watch the 50M dollar cardiovascular grant program for whether it publishes actual safety endpoints and whether those endpoints are defined rigorously or defined to be passed. A program designed to generate the data that justifies a conclusion already reached is not research, it is a procurement exercise with a literature review attached. The quality of the endpoints will reveal which one this is. Watch the malpractice

carriers, because the first carrier to write a policy for an autonomous AI prescriber, and the terms they write it on, will reveal how the people whose actual money is at risk are pricing the danger that the founders wave off.

Watch the litigation. The Character.AI action in Pennsylvania is the opening shot, and the unauthorized practice of medicine theory is the one most likely to scale across states, because it does not require new law, it just requires applying statutes that already exist. If a state AG wins a clean case on that theory, it changes the calculus for every company in the space overnight.

And watch the failure that has not happened yet, because it will. Wachter is almost certainly right that at some point a system gets trusted past what it has earned and someone dies, and the question that determines the next decade is not whether that happens but what the regulatory response looks like when it does. If it triggers a hard rollback, autonomy gets set back years. If it gets absorbed as an acceptable cost of progress, the way a certain number of robotaxi incidents have been absorbed, then the genie really does stay out of the bottle. The uncomfortable truth sitting under the whole debate is that American healthcare already kills a lot of people through ordinary human error, and the strongest argument the optimists have is not that AI is safe, it is that the bar it has to clear is lower than anyone wants to admit. Whether that argument should win is a values question dressed up as a technical one, and the people building the on-ramp are betting the country answers it in their favor before it fully understands the question.

support my work, consider becoming a
free or paid subscriber.

← Previous

Discussion about this post

Comments Restacks



Write a comment...