



## Podcast Transcript

# Regulatory Horizons: AI & Digital Health – FDA Evolving Oversight

**Date:** January 6, 2026

**Guest:** Winston Kirton, James Sherer, Aleksandra Vold, **Host:** Amy Kattman

**Run Time:** 35:12

**For questions and comments contact:**



**Winston S. Kirton**

Partner  
Washington, D.C.  
T: 202.861.1715 | [wkirton@bakerlaw.com](mailto:wkirton@bakerlaw.com)



**James A. Sherer**

Partner  
New York  
T: 212.589.4279 | [jsherer@bakerlaw.com](mailto:jsherer@bakerlaw.com)



**Aleksandra Vold**

Partner  
Chicago  
T: 312.416.6249 | [avold@bakerlaw.com](mailto:avold@bakerlaw.com)

---

Kattman: Welcome to Regulatory Horizons, delivering strategic insights for the future of global life sciences regulation. This podcast is your guide to the evolving policies, global frameworks, and innovations shaping the compliance landscape across

biotech, pharma, medical devices, and even foods. I'm Amy Kattman, and you're listening to BakerHosts.

Today, we will focus our discussion on AI and digital health and the FDA's evolving oversight. The FDA has been actively shaping its regulatory framework for AI and machine learning in medical devices, recognizing that traditional device regulations were designed for static products, not adaptive algorithms that learn over time. Let's listen in.

Kirton: Welcome to Regulatory Horizons, delivering strategic insights for the future of global life sciences regulation. This podcast is your guide to the evolving policies, global frameworks, and innovations shaping the compliance landscape across biotech, pharma, medical devices, and even foods. My name is Winston Kirton, and I am the leader of BakerHostetler's FDA practice and co-leader of the firm's Life Sciences Industry team.

Today, we will focus our discussion on AI and digital health and FDA's evolving oversight. The FDA has been actively shaping its regulatory framework for AI and machine learning in medical devices. FDA recognizes that traditional device regulations were designed for static products, not adaptive algorithms that learn over time.

There are three key principles in the associated framework that FDA is working off of. The first is software as a medical device. So, AI-enabled software intended for diagnosis, treatment, or prevention of disease falls under the software as a medical device category. The FDA aligns its approach with the International Medical Device Regulators Forum principles, which emphasize clinical evaluation, transparency, and lifecycle oversight. The second key principle and associated framework is total product lifecycle approach. So, oversight spans from development through post-marketing monitoring, including green learning practices and the guidelines for data quality algorithm validation and bias mitigation. The last key principle is predetermined change control plans, and these allow manufacturers to predefine how an AI algorithm can evolve post-approval without requiring a new submission for every update critical for adaptive AI systems.

Now, FDA has put out guidance, three key guidance documents, the most recent being in January 2025, which was entitled AI-Enabled Device Software Functions: Lifecycle Management and Marketing Submission Recommendations. Prior to that, in December 2024, the agency put out Marketing Submission Recommendations for Predetermined Change Control Plans. And in 2023, the agency put out Cybersecurity in Medical Devices guidance. Now, most AI machine learning-enabled devices are cleared through the 510(k) process, in fact, about 96% of AI ML-enabled devices. So while high-risk or novel technologies go through PMA or the De Novo classification route, overwhelmingly, these enabled devices are cleared through the 510(k) process.

Now, what are the challenges for the life sciences industry? One is transparency and bias, so FDA emphasizes explainability and bias mitigation in AI models. The

next is post-marketing surveillance, so continuous monitoring for algorithm drift and performance degradation is now a regulatory expectation. And lastly, cybersecurity and data integrity. AI systems must safeguard patient data and maintain reliability under real-world conditions.

So hopefully, the audience has kept this frame in mind, and I ask that you keep it in mind today as I'm pleased to be joined by two of my esteemed BakerHostetler partners who are leading experts in AI and digital health, James Sherer and Aleksandra Vold. So James co-leads the emerging technology team of BakerHostetler's Digital Assets and Data Management Group, while he also directs the firm's artificial intelligence and information governance engagements. His works span litigation, transactions, and regulatory enforcement, while helping clients confront discovery management process issues, enterprise risk management, records and information governance, data privacy, security and bank secrecy, artificial intelligence and algorithmic transparency, technology-integrated issues, and related merger and acquisition asset purchase and divestiture diligence. Aleks Vold has handled hundreds of healthcare breaches and has played a lead role in some of the largest breaches in history. Aleks is the firm's Chicago Digital Assets and Data Management leader, and she frequently advises clients on compliance with state, federal, and international law and regulations. James and Aleks, welcome to the Regulatory Horizons podcast.

Vold: Thank you, Winston.

Kirton: Yeah. So let's get to it. I was hoping that both of you can share how you became involved in AI and digital health regulation for our audience.

Vold: Sure. I can, I'll go first. So in my role as HIPAA counsel, as a counsel for large healthcare systems and business associates alike, as this wash of AI innovation has come on the scene in the last several years, it went from being a low, an emerging issue to one that is being brought to the privacy offices, which is who I deal with. And so starting a few years ago, we went from normal discussions around, is this a good idea? Can we do this under HIPAA with respect to normal SaaS product providers? To, hey, we are interested in or we've got an initiative from a particular physician around wanting to use AI for these new use cases, and we're just wondering if the rubric, the business associate rubric, the HIPAA rubric is really the same and what we should be additionally calling out, if anything. And so it was a bit by just through questions of clients that have grown and grown over the years to be more prevalent and more top of mind. And so, fairly organically from just normal day-to-day HIPAA privacy and security role advising has become a significant part of my practice.

Kirton: And how about you, James? Thank you, Aleks.

Sherer: So I want to take you back to the, I suppose the early 2010s when AI was branded big data, and that was the term du jour and what the marketing professionals were using. We had done some work into advanced data analytics and looking at larger data pools and environments where these cross-company analytics could be performed. And we were approached by, really, a venture

within the healthcare space that sought to combine covered life data among, I think it was 37 companies to start with, that neared, I think, 10 million covered lives, where those entities wanted to combine the data from PBMs and providers and insurance companies into one place to see if they could improve course of care practices.

So it wasn't called AI, but that was also when you saw the emergence of the IBM Watson platform, looking at what Truven was doing in the space, and that led to really refining what approaches would look like for the application of AI in the healthcare space. And those ended up being pretty unique because while there are a number of industries where you don't want to make distinctions based on certain characteristics of people that would otherwise create disparate impact, within healthcare, it really is so individual focused. And we want to leave opportunity and availability to analyze the nuance of people, every patient being that snowflake, but also looking at data as a whole. Looking at trends and implications where, especially for the first initiatives, even these very minuscule or marginal improvements in care and cost, when you advertise them or spread them across larger groups, they really can be impactful. And if we have an opportunity, and looking at like client approaches, to improve care and also save money, which, by extension, if you're saving money and plugging it back into the system to therefore improve care, you can see a little bit of a snowball effect. And that was and continues to be, I think, the primary focus of AI within healthcare.

Kirton: Staying with you, James, what excites you the most about the potential of AI in healthcare and life sciences?

Sherer: That there are some objective good benefits that we can see. Unlike the application of some of these systems where maybe it is how, if we can best sell a sweater or if we're looking at improvements on automated marketing, we can look at the application of these systems where we see unmitigated good. We see improvements in courses of care and patient outcomes, and really the industry embracing that.

One of the seminal examples has been, we see improvements in reading outputs from systems where suddenly these AI systems are looking at it with these different facets and coming up with differential diagnoses, giving a little bit more depth to approach, catching things earlier or that just would have otherwise escaped detection from practitioners, and then improving patient outcome. And it is not one of those things where, unlike some of the other industries where the professionals are saying, oh, wait, AI is going to steal my job. Here, you've got medical professionals saying, this is wonderful. This is improvement to patients. What can we learn from these systems and is it translatable even into what we use in our education process? And sometimes it is not, which means we are moving into this area of contributory medicine. We're working hand in glove with the technologies, they really are augmenting what professionals are doing. And it should, and the hope is, of course, that it is going to democratize a lot of these areas of care. It is going to improve care. If we can scale good medicinal practice, device use, improve medicines across populations, that is good. No one is debating whether or not improving care and outcomes and maybe alleviating

suffering and extending life and causing quality of life to increase is a good thing. So I think that is where everyone is rowing in the same direction and I think that is where the excitement comes from. And it is exciting to help support that as a legal professional.

Kirton: Wonderful. Aleks, I want to hear a little bit about your excitement as well, if you wouldn't mind sharing with listeners.

Vold: Yeah. I think that there is the federation of knowledge gained by AI, so it is compounding excitement, right? You're compounding returns here, right? So each individual deployment of AI has the ability to much more quickly identify patterns, successes, failures, et cetera. And when connected in the right community, the federation of those learnings, so we can share learnings without sharing data or compromising patient privacy, means that everyone can benefit from this. And we're all singing off the same sheet of music as well, everyone being comfortable with the systems in place, how these learnings were identified, and so on and so forth. And I think it has the ability to break down silos in research much more quickly and easily. And then again, to James's point, bettering the patient experience.

And along on the administrative side, we've heard about physician burnout for a long time, certainly. And not just physicians but any bedside, that was definitely increased through the pandemic and having this potential hole in the caregiving space. And so anything on the administrative side could ease some of that burnout in a safe way, in a way that is reliable, is fantastic. We talk about, I was at a roundtable with some healthcare executives a couple of weeks ago, and the CEO of a hospital system was talking about AI increasing pajama time, that time with the physicians with their kids, making sure they can be home for that instead of being at work charting and doing their dictation and notes. And although that is not a direct monetary gain necessarily, in the long run, keeping these folks engaged, keeping these physicians and nurses and caretakers in the industry has extraordinary benefits, not only financially, but also from a patient care perspective and continuity of care. Losing these folks is a big problem that a lot of our clients are facing and trying to figure out how to do this, to continue to staff and staff well. So, loving to see that in addition to the clinical applications.

Kirton: Yeah, no, it is wonderful. I think we, as a firm, continue to see an uptick, the education around AI, wanting, clients wanting education on AI and clients venturing into AI. And of course, we were able to provide that one-stop shopping, if you will. I'd like to stay with you, Aleks, and transition to our clients, companies that we serve, and get your point of view on what are some of the biggest hurdles you're seeing our clients or companies in general in the healthcare life sciences space. What are some of the hurdles they're facing when they contemplate bringing AI-enabled products to the market?

Vold: Yeah. So I think one of the biggest hurdles is, so AI, there already being a backbone, right? You can incorporate already created AI tools and companies can use those tools in unique ways. So barrier to entry, I feel like, is lower than if you're building an app from scratch, right? And so I think there is a deluge of new

vendors catering to the medical field, saying things like, we're going to disrupt healthcare, which love the energy, but we're a highly regulated industry, highly regulated industry. And that is not always appreciated by these new vendors or even vendors who've been around a long time that want to incorporate new AI tooling in longstanding pieces of software. And so I think the biggest hurdles are, especially for new entrants into healthcare, understanding that AI in retail and AI in healthcare is different. We have different privacy issues that we need to deal with. We have different applications, and we can't just break stuff and go fast in healthcare a lot of times. We can try to go fast, but we certainly can't break HIPAA just to get this AI-enabled product in our systems, no matter what the great result is, because we're going to eventually be on the hook as a covered entity. So getting them to understand those sensitivities and the realities of being in the healthcare market are sometimes a big hurdle.

And then the negotiation process again, because of these new entrants, we're seeing terms that we never see from vendors who have been in this space a long time. The other day, I saw a business associate agreement that said that this AI vendor was only going to be liable for breaches that happen in the first 12 months of the contract. That makes no sense at all, it is not tied to reality. I'm not sure why, but that is a clear showing of immaturity in that vendor itself.

So it is really just marrying how can we do this with HIPAA? And it is absolutely possible to deploy AI in a way that is HIPAA compliant, and we'll get to a yes. It is just, it takes some poking and prodding with the vendor and with the software and with the champions within this system about, how is this going to be used, how can we limit, how are we making sure that we limit the vendor's use of this data? Because I think that is another issue here is when you're built on AI, the idea is always, of course, to make the AI better and use what your customers are providing to you to better everyone's experience. But that is not necessarily, I'm not going to say across the board it is not okay under HIPAA. But in some situations, depending on exactly the terms and the data, that may not be acceptable, or at least not an acceptable risk to the covered entity. So ensuring that you're really keeping an eye on the secondary use of data is another bit of heartburn that we see on the covered entity side.

Kirton: Thanks. Thanks, Aleks. James, I want to switch to you. Aleks mentioned counseling clients that, wait a minute, this is a very regulated environment that you're entering into. And we do see a fair amount of, say, startup companies, if you will, as well as mature companies who have resources. But what advice would you give to our listeners that are in either category, either startup or mature company, about ensuring compliance early in the development of AI products for healthcare and life sciences?

Sherer: Because we've already said this might be attorney advertising, should come speak with us. But given that, and it would be the same context regardless for that conversation, it harkens back to the data governance approach and some of those sensibilities, where a common saying there is you don't have to boil the ocean. We want to look carefully at where to deploy limited resources.

The FDA guidance that we've seen, as well as normal industry practice and general approach to AI, speaks to a risk-adjusted approach. Looking at what you have, looking at where it is being deployed, thinking carefully about the ramifications of different actions, and then trying to orient those things toward the art of the practicable. What's reality look like there? And given that the FDA is still waiting, as I understand it, or to complete the feedback loop for some of the requests they have out there in the industry, and obviously different approaches given different administrations, I think I tend to draw those conversations back to what we see as common themes.

So within the United States, within Europe, we've seen direction back to emphasis on documentation, model cards, understanding, obviously, how these technologies operate. And then looking at, depending on how you articulate it, explainability, interpretability, transparency, which I don't love, looking at the pure play math of this, may not be helpful to someone who is considering deploying or even utilizing a technology. But there are variations on the same theme of building out a program, asking pretty cogent questions at the beginnings of it, and then understanding, to a point that Aleks made, understand the trajectory of your technology. Where is it going to intersect with the market? What expectations are going to follow that if you're the developer as opposed to the deployer of the technology? If you're a deployer, are you adding different technologies together and creating something new? What kind of obligations does that attach then?

One company that I work with put it most poignantly. They said, it is like we need an HR for AI. If we've got this idea that AI is acting as a person but it is scaling, if it is going to help with that bedside manner and help with a different physician or medical practitioner burnout, we have to appreciate the counter to that, which is if it approaches things differently, if a mistake is made, is that mistake going to scale?

So when I work with clients, and I've worked with Aleks on these points as well, I think when in doubt, I try to orient them back toward considerations of different frameworks and how those might apply. And it is not to say you have to do everything. Again, don't boil that proverbial ocean. But what is the art of the practical? We look at things like the NIST AI Risk Management Framework or the ISO/IEC 42001:2023. Both of those frameworks have been recognized as extending safe harbor protections under upcoming U.S. law related specifically to AI and were parts of considerations for things like the EU AI Act. And even the FDA guidance looks to international implications because these technologies, especially when they can better courses of care and improve healthcare outcomes, are going to ignore geographies. And I think and expect that humanity would approve of that and support it.

So long story long, to that question, when in doubt, I try to work through things as benign as checklists and say, here is overall what we would like to see. As we answer some of these questions, we look at controls and compensating controls associated with AI tool development and deployment. Let's see what we have. Let's see what questions we can answer. And the journey of evaluating those

questions, looking at policies, looking at model cards, understanding the data provenance of what is going into these systems, what they're built upon, and then where they're extending out into the market, all of those questions are extraordinarily important to understand what the risk is that the organization might be confronting. And then it'll help to deploy often limited resources in the ways that are going to make the most sense for a given organization.

Kirton: Yeah. I learn something new every time I host one of these podcasts, right? And today I learned HR for AI. So James, you might have to come on a podcast and just explain that one. Sounds simple, but we get the point.

Aleks, you mentioned patient protection earlier, and as one of the challenges for both healthcare and life sciences is around cybersecurity and data integrity. We're seeing an uptick in cybersecurity incidents impacting life sciences companies who are in the manufacturing space and doing clinical research and the like. And so we are now working with companies to ensure that they're meeting their regulatory obligations around cybersecurity data protection. But knowing that you live in this data protection world in many instances, can you tell us one or two approaches that you've found that have been successful for companies in regards to cybersecurity and patient data protection in AI systems?

Vold: Yeah. I think the most important thing to remember is that PHI, protected health information under HIPAA, remains PHI unless it is truly de-identified. And that either means the removal of 18 specific identifiers as identified in HIPAA, which basically renders the data almost useless, and it is things, it is, you could have nothing else about me other than my date of admission, and that would be not de-identified PHI if there is any other clinical information. So, not de-identified just by taking my name out or my address out. There is a lot that goes in that or getting an expert opinion. And so I think that the thought a lot of times is that, oh, this \_\_\_\_\_ doesn't say a patient name or it doesn't have the MRN. And so it is no longer PHI and where, if there is a breach, it is no big deal. It is big deal. It is definitely problematic and still subject to HIPAA unless you've got that de-identification down pat.

And HIPAA is, in the sense that it tells you very specifically the kinds of controls you need to have in place, the logging that needs to be done, that it needs to be reviewed. But treating repositories, the prompt injections, whatever you want to call them, that the tool is using, treating that very carefully and protecting that in a way that you would if it was more explicit PHI than any other system. A lot of that, a lot of the issues that we see in data breaches are related to human error, whether it is someone clicked a button to make something public that was not supposed to be public, they've hooked their credentials when they thought it was a legitimate e-mail that allows for access remotely to their network, doesn't patch their server in time like they were supposed to, whatever it is, these are all human error issues. And so ensuring that you've got policy, the policy itself is not going to protect PHI, but enforcing them, ensuring that you're enforcing your security policies is absolutely crucial because you're never going to stop people from clicking the wrong link or entering the passwords. And so then it becomes a matter of early detection and making it as hard as possible to get at that data and

then seeing the logs to see what happened if slash when a bad guy got into your data repository.

The other issue to think about is AI is being used by threat actors too, right? So these are no longer just some kid typing away on their computer trying to one password at a time, it hasn't been that way for a long time, but one password at a time trying to get in. Now, they are getting in through the social engineering and then using AI to better understand system topography, find data, reduce protections that are in place. And so constantly thinking through whether your protections are up to snuff, whether you can do more within reason, it is going to become even more important because the threat actors have always been one step ahead. That is why we have breaches for a lot of companies. But that pace is going to continue to exceed the pace of protections because we have budgets and need to get approval and need to deploy, whereas they are just running full steam ahead with these tool sets. So protect the data as if it were full PHI, even if you're not thinking that way, and that is going to get you at least in a good position to spot things happening before it gets too far.

Kirton: Thank you so much for that, Aleks. One of the benefits of working at Baker is that across practices our lawyers tend to focus on business and helping clients manage through business, even, or in spite of, regulatory constrictions and regulatory, deep regulatory environments. We like to, in this podcast, we like to share information, we like to educate, and then we like to, at the end of it, wrap it all up with inspiration.

So, I want to stay with you and then go over to James in the last couple minutes we have. I want to talk a bit, get your feedback on innovation and where you see the biggest opportunities for innovation with AI under the current regulatory frameworks that you've been exposed to, including, we have mentioned, of course, FDA, and of course, there is other regulators. So what is the biggest opportunities for innovation? And then, if you have some more insight, we'd love for you to weave in, what would regulators have to do to evolve as well? So evolution and innovation in products and the regulatory environment, but also regulatory evolution. We'll stay with you, Aleks, and then end with James.

Vold: That is a very big question to end with.

Kirton: It is, right?

Vold: It is a big question. I truly think that the opportunities are so expansive. And particularly in the research space, there is so much here that it is not taking away jobs. It is not being flippant with we want to use patient data to advertise better. Or even if you want to call pajama time insignificant, which I don't, but if that is not hitting for any of the listeners, I think on the research side, one, we have a lot more flexibility under HIPAA in sharing data for research purposes. And even under the more restrictive Part 2 for substance use disorder, that is another more restrictive space on patient data sharing, both HIPAA and Part 2 allow for different types and more open sharing of patient information sometimes. And I think using that and running with it will not only have significant impacts on

patient care and care modalities and even the life cycle of what it looks like to be a physician, having answers faster and that are more well-informed is just, I think, so exciting.

And I do think that there is a difficult path forward because we've got a lot of, we've got states coming online in the next year with new AI laws. We've got federal government saying we're coming for this either from an HHS OCR perspective, the White House generally, we have the FDA. There is going to be, I think, the real challenge is getting everyone on the same page and not creating such a patchwork that it becomes unnavigable. And so I think some of the opportunities here don't necessarily lie with the implementer, the creator of the AI, but with our various government officials and deciding, what is the most important thing here? How do we implement something that does not create risk for patients and their data, but that can move the country forward and healthcare forward in a meaningful way? And that government is slow, regulations do not change overnight. But we really do have an opportunity here to be both protective and innovative if we can get the stakeholders on the government side behind a single goal here.

Kirton: Wonderful. Aspirational. Inspirational. James?

Sherer: So I think mine is more practical. This is less about the hope and anticipation, and it is more just the practical issue of it. We look at the novelty of some of these technologies and certainly generative AI and how that captured, I think, marketing and imagination of people, making it accessible to individuals rather than requiring entire systems and tech stacks to develop things that now Gen AI can do. So we see that and we see that material on the surface, but a lot of the laws that are getting implemented are not just encompassing these new technologies but also incorporating sensibilities from things before.

I'd started off by saying some practitioners look at AI as just being big data rebranded. This is an extension out of evaluating these large sets of data and coming up with new applications and new insights and then implementing those. So, for organizations, I think it is a nice return to good foundational documentation and understanding. That is where we're going to go. It is less a matter of all of the hope, which I want to leave out there, but you get good results from good foundations. And that is where I try to orient clients. It is not just, let's chase the next thing that is shiny, but let's understand what we have. And appreciate that this may be a real turning point for that, because organizations are not only going to be tasked with understanding better the technologies they already have in place, but there is a lot of pressures for organizations either to retool their workforces to take more advantage of these technologies in order to reposition themselves to see where the market is going to go.

So there is still a lot of opportunity, and I would assert a requirement for common sense in looking at these things and saying, we already know how to handle some technologies. AI is no different for some of the foundational questions, and we're going to need to accommodate for that, even while understanding there

may be some novel uses that we're also going to have to consider once we have our feet firmly set on the ground.

Kirton: We're out of time and clear, practical, actionable advice. And James and Aleks, I really appreciate you joining the Regulatory Horizons podcast. And I hope that we can invite you back either together or individually as the regulators continue to build out their frameworks and companies continue to work on ways to advance AI for their specific products in life sciences and for the services as well as products in healthcare. So thank you. We thank our audience for listening. Stay tuned for the next in our series in the Regulatory Horizons podcast. Thank you and be well.

Kattman: Thank you, Winston, James, and Aleks. If you have any questions for Winston, James, or Aleks, their contact information is in the show notes. As always, thanks for listening to BakerHosts.

Comments heard on BakerHosts are for informational purposes and should not be construed as legal advice regarding any specific facts or circumstances. Listeners should not act upon the information provided on BakerHosts without first consulting with a lawyer directly. The opinions expressed on BakerHosts are those of participants appearing on the program and do not necessarily reflect those of the firm. For more information about our practices and experience, please visit [bakerlaw.com](http://bakerlaw.com).