

[music]

Geri Amori, PhD, ARM, DFASHRM, CPHRM: Hello everyone, and welcome to *Healthcare Perspectives 360*, a podcast dedicated to exploring contemporary healthcare issues from multiple perspectives. I'm Geri Amori, and today I'm joined by Irene Dankwa-Mullan, MD, MPH, a physician, executive researcher, and thought leader in health technology innovation, currently adjunct professor at George Washington University's Milken Institute School of Public Health, Chief Health Officer at Marti Health, with expertise on the ethical use of AI in healthcare, aiming to bridge gaps in access for underserved communities and ensure personalized, precision care for equitable outcomes.

I'm also joined by Danielle Bitterman, MD, assistant professor at Harvard Medical School, and she is also a radiation oncologist at Dana-Farber and Mass General Brigham, and she has unique expertise in AI applications for cancer and AI oversight for healthcare. Also with us today is Chad Brouillard, Esq., a medical malpractice defense attorney and a partner at the law firm of Foster and Eldridge, LLP, in Massachusetts, with expertise in electronic health record liability, AI implications, and all things technologically related to healthcare, liability, and risk. Welcome to our panelists and welcome to our audience.

Today's theme is strategies to address the risks of AI for patients and providers: how can AI be safer? So I'd like to start today with you, Chad. Something I learned about you recently is that you are not only an attorney, but you have studied philosophy and ethics. That's an amazing combination, a lawyer and ethics. Okay. So I'd like to kick off today's session by asking you a general question. When looking at the implementation of new programs, offerings, or methods of computerized support, what are some of the general questions that we in healthcare should be asking, like, how do we move into an AI support ethically?

Chad Brouillard, Esq.: That's a really great question, and a noted observation that it's a contradiction, right? An ethical lawyer, imagine that. I think the overarching 10,000-foot view is, you know, just because we can implement technology, does that always mean we should, right? You know, the mantra of the Silicon Valley is, move fast and break things. But that should give us pause, right? When we're talking about cutting-edge technology in the healthcare space. I think, you know, on one hand, there are forms of AI that have been around for decades. 1995 is the first application that's approved by the FDA, for instance, right? By other hand, the current generation of large language models, generative AI, you know, there is an element of almost experimental in the medical space. And so I think it requires a lot of things to advance this in an ethical way. I think from due diligence point of view, before implementing technology like this, really vetting it pre-adoption is such a core function, and it really needs to be done sort of in a multidisciplinary sense because you have so many structures in play.

So, you know, obviously clinicians, your health information technology departments, yeah, of course they need to be involved. But sometimes I think systems don't think of risk managers, right, who know rules and know how implementing technology in a certain way might actually be breaking more traditional rules, things like making sure you have continuous testing of outputs and the ability to give feedback to the vendors, thinking through things like are there times where we need to get patient consent to use this technology, and how does that process look like? And is there a process for opting out of the use of AI at a particular facility or department based on the tool?

I think overall what's needed in terms of AI implementation in the healthcare space is careful monitoring from a multidisciplinary perspective and information governance committee that is really looking at the tools in place on a monthly or even more basis, giving feedback from various perspectives and making sure that there's good testing from these applications.

Amori: Okay, so move slowly and don't break things, basically.

Brouillard: You got it.

Amori: So Irene, I also know that you are involved in ethical processes and have a great interest in this. So, what types of decision tools or processes do you feel are central for ethical implementation of AI?

Irene Dankwa-Mullan, MD, MPH: Yeah, thanks. That's an excellent question, and it actually gets right to the heart of ensuring AI in healthcare doesn't just advance technically but also ethically. And so I've published in the space, and there are few key decision tools and processes that I believe are central to the ethical implementation of AI in healthcare, and the first is bias-auditing tools. I think they are crucial for identifying and addressing biases in AI models. We need robust auditing and validation frameworks because we know AI systems are only as good as the data that they're trained on, and it's especially crucial or critical for populations that are often underrepresented in training data sets.

The second is explainability. AI systems need to provide clear, understandable reasons for their recommendations, especially in clinical settings where trust is paramount, like we often talk about black box. Black box models that deliver predictions without explanations can lead to hesitancy and even rejection by clinicians. And so that's one and then thirdly, we really need governance frameworks that focuses on accountability and oversight. This means that clear policies about who is responsible when the AI gets wrong, how organizations/hospitals need to track these tools. Governance should also ensure that the patients have a say in how AI is being used in their care. I think, you know, Chad had mentioned that, which ties to the principle of patient autonomy. Fourth is equitable data practices. Really essential, right? It's not just about having a lot of data. It's about having the right quality data. It's about making sure that your training data sets have diverse population.

A really important one that we really don't consider but needs to really emphasize is continuous monitoring and feedback loops. AI shouldn't be just set and forget; there needs to be continuous evaluation, monitoring the impact, identifying and adjusting the algorithms, accordingly, making sure that these tools really remain relevant and accurate and fair over time.

Amori: So you're already, you're giving us some great safeguards and ways to make sure it continues to have ethical use. Danielle, in light of what's in place now, do you see areas that need additional attention surrounding the ethical AI decision-making process?

Danielle Bitterman, MD: Yeah. And I asked first of all, just echo that Chad said, move fast and break things does not work in healthcare. It is a recipe to get clinicians and patients to rapidly lose their confidence in AI, and that's actually going to set us back. But in terms of where additional work is needed now comes down to the fact that right now a lot of the attention is going to the models themselves, but the models are actually just one small piece of the full implementation the software's medical device that those models will live within, how the outputs of the model is communicated to the

end user, the patient or the clinician, how it connects to other data pieces, oftentimes the more challenging, complex portion. And I think there's a lot of room to improve in terms of embedding ethics by design and safety by design into those processes, making it when you report something of a language model, for example, potentially providing some brief information and easy to digest education on where those will be found. Those models to sometimes have errors. What you should pay particular attention to can really help, for example, clinicians be effective overseers of those models. So while that's not particularly kind of focusing on optimizing the model itself, it is an area that's not getting quite as much of attention as it should but really is needed to make sure that we stay vigilant are good and effective and responsible users of the AI that's integrated into our systems.

Amori: Okay, all right. So Chad, I'm going to skip to you and ask you a question here. One question that has come up a lot is the need for patient consent when any form of AI is being used. And the most conservative arguments are that patients should be told about how medical recommendations are being made, including any collaboration with AI-generated literature, which I guess there's lots of that – I know we all do searches all the time, right – or screening tools. So at the other end of this argument is that we don't tell patients when we use spellcheck or some other form of AI, so using your lawyer language, it doesn't feel like there's a bright line, but it also feels like a really big slippery slope. Help. What do we need to consider when we think about where's the line, the place where we need patient consent?

Brouillard: Again, to quote my law school professor, depends, right? So you can take the consent piece and really spin it out to absurdity as you suggested, right? You don't need to consent a patient every time you use a spell check function, right? But I do think there are some sources suggesting the development of lines of where consent is probably needed. I think the World Health Organization has ethical guidelines, at least for suggestion, and they kind of tease out a thought problem about this. But the suggestion is, generally, anytime AI is being used as a substitute or augment for clinical judgment, and if the AI is introducing any element of risk of inaccuracy, probably that needs to be addressed. Now, part of the problem is, if we are embedding a lot of different AI tools, that's going to be different. Do I have to keep stopping and asking, Can I use this tool now? Can I use this tool now? I mean, I think logistically, I think a lot of healthcare systems are building it right into their consents of treatment and trying to explain to patients that they might use a variety. I think it's very different if there is a particular intervention, such as a surgery, and you're using AI, then there's a whole consent, formal consent process that's in place already to begin with.

I do think that the other big consent, putting aside its ability to introduce risk and the need to inform patients about that risk, I do think there is sort of a consent process in terms of using, retaining and repurposing the patient's data, right, whether it's for research purposes or clinical purposes, particularly if you're using third-party vendors. I think that's the other bright line is that it may require you to get consent if, you know, for instance, that by using this AI product, you're having to store patient data outside of the facility or share it, or they're going to repurpose it, then that might be a point where this is needed. I also just want to point out this can be very, very difficult for clinicians being aware that AI is even being used, particularly in the context of a medical tool or the electronic health record. And I think that raises another difficulty of even knowing, the challenge of even knowing, sometimes, when these tools are being deployed, when the information is being presented on the clinician's screen.

Amori: Okay, all right. Well, Danielle, I guess I want to ask you now. We've talked about whether or not we need to tell the patient, but we also know that AI models don't go through the rigorous trials say that new medications go through or new treatment processes go through. So on what basis should we trust them and ask patients to trust them? What should be the guideline on that?

Bitterman: So Chad's turning me into a lawyer, so I would say it depends to this.

Amori: This is going to be the "it depends" session.

Bitterman: So not every new intervention or new software as medical device requires a full, randomized clinical trial. I'll put that kind of to lay the groundwork. The level of evidence you need to adopt something in clinical care depends on the risk to the patient and kind of whether or not how it affects patients' other treatment options. We don't necessarily need a full, randomized clinical trial for every AI implementation, and that's also not practical. At the same time AI is advancing – even ones that you would normally like to do a full clinical trial – the methods are advancing so rapidly sometimes by time you finish the trial, you're now kind of generations ahead in terms of the new technologies, and you don't know if the results of the original trial are valid.

So, I think there's a lot of room for newer trial designs, pragmatic trial design, so more nimbly test these. So it's hard with trust. You kind of have to trust something in order to trust it. So, I think gathering the appropriate level of evidence for that given intended use of your system is an important basis for trust. Understanding what is a measurable endpoint you want to say that this is helpful or harmful is important to communicate to patients and clinicians for trust. And for patients, being upfront with telling them when we do not have full evidence. We do not, this is an experimental setting, and clearly explaining benefits and risk is standard and a reasonable way to communicate with patients, and I think actually essential for having them trust us, and as we learn how to safely use these models.

Amori: Okay, so learning the basis for trust. Well, Irene, that leads me to you, and partially because, you know, we're sitting here and we're very skeptical and we analyze, but some people, actually, are prone to trust computers more than people. They just are. I mean, I have a grandchild who says, Well, this thing site says this person is a doctor, and I'm known for saying to her, just because it says they have an MD doesn't mean they're smart, it means they went to medical school, right? So, how do we figure out a way so that AI doesn't erode patients' trust in legitimate healthcare providers? What about those deepfake providers on the media? How are we going to work with that?

Dankwa-Mullan: I know that's a really important concern, very concerning, especially, I mean, patients might question their clinician or their doctor's recommendation if it differs further from what AI suggests to them. So it's a real, real challenge, and of course, the rise of deepfake providers only makes it more complicated because they can really present false information. So I think the solution lies in transparency and education. I think AI tools need to disclose how their recommendations are generated, what data was used. I'm always looking for what data was used, what were your, you know, what did these sort of use as your features right to make those decisions. It's really important and any limitations that they have, right? And I think as healthcare providers or clinicians, we need to be proactive in explaining the role of AI in care decisions, making it clear that AI is only a support tool. It's not a replacement for an expertise. So really stronger regulations AI detection tools, but in short, building trust will really require a balance of education, transparency, stronger safeguards against misinformation.

Amori: And maybe teaching some of our patients a little skepticalness, too, doing the things they find on the internet. Yeah.

So we've come to the point where I need to ask today's, my favorite question. If you want our audience to take away one thought from today, what would that be? And Chad, I'd like to start with you, if you don't mind.

Brouillard: Sure, Geri. I mean, I think overall, and I agree with everything that both Irene and Danielle had to say today. I mean, I think the thing is that we all agree that AI solutions in the clinical space really need to be rolled out in a very thoughtful and careful manner. And you know, for my piece, I would say that embedding it in an organization's Information Governance Program is just so essential, so that you have continuing eyes on the outputs over time, and you've done proper vetting on the front end, you know, to whatever degree you have a consensus about, right? And I just think that's so important. So, the solutions have great promise, but they have to be implemented carefully.

Amori: Okay, promise, but implemented carefully. Irene, what would you like to say one thing for our audience to take away today?

Dankwa-Mullan: I mean, clinicians are concerned about the safety and accuracy and transparency of AI recommendations, especially when these recommendations influence sort of high-stakes decisions without clear explanations. And so the bottom line is, I think for AI to be safer for both patients and for healthcare providers, we need that foundation of transparency, robust bias audits and clear accountability. When AI is transparent and accountability, it has the potential to enhance, right, not necessarily please the trust and expertise at the core of effective healthcare.

Amori: Good. That's a really important point. And Danielle, what would you like our audience to take away from today?

Bitterman: I mean, I would say that AI is moving fast, and there's a lot of new tools that we're going to start seeing. But at the end of the day, the core bioethical principles that we in medicine adhere to, that are the standard for guiding how we conduct research in medicine, how we deliver clinical care, aren't changing. We don't have to remake the wheel. We want to put good-faith effort into evaluating and monitoring and implementing these new technologies so that they adhere to existing ethical principles and protect patients' rights, so we have a long-standing basis that we can learn from and an infrastructure that we can build on top of to do this right.

Amori: Excellent. That's a really, really good point. I just want to thank our panelists. This has been a very, very interesting conversation. We think of AI and all the technology, but there's a whole human aspect, the ethics and the application and the approach. So I just want to thank you, and I'd like to thank our audience for participating today. And I will see you again, I hope, next time when we explore another medical issue from a *Perspective 360*.

[music]