

Transcript Details

This is a transcript of an educational program. Details about the program and additional media formats for the program are accessible by visiting: <https://reachmd.com/programs/the-convergence/promise-peril-clinical-ai/54783/>

ReachMD

www.reachmd.com
info@reachmd.com
(866) 423-7849

The Promise and Peril of Clinical AI

Dr. McDonough:

Clinical AI is moving from a tool that answers questions to something closer to an agent that can help coordinate care, reach patients between shifts, and close clinical loops. That shift raises a practical question for clinicians: when is AI simply helping us think, and when is it acting on behalf of the care team?

Welcome to *The Convergence* on ReachMD. I'm Dr. Brian McDonough, and joining me to explore this question are two returning guests: Dr. Raja-Elie Abdulnour and Jean-Claude Saghbini.

Raja-Elie is an ICU physician, Editor-in-Chief of NEJM Clinician, Chief Clinical Innovation Officer of NEJM Group, and Assistant Professor of Medicine at Harvard Medical School. Raja-Elie, welcome back to the program.

Dr. Abdulnour:

It's good to be back, Brian.

Dr. McDonough:

Jean-Claude is Chief Technology Officer and President of Technology Solutions at Lumeris, where he leads the development of AI-powered solutions and their integration into clinical care. Jean-Claude, it's great to have you back as well.

Mr. Saghbini:

Thanks for having me here.

Dr. McDonough:

Raja-Elie, in our prior conversation, you described AI in high-stakes care as something like a very good resident—useful and knowledgeable but still requiring supervision. Jean-Claude, Tom is designed to act on behalf of the care team between visits. So let's start with the core question: when should AI be allowed to simply advise, and when should it be allowed to act?

Dr. Abdulnour:

I'll take that question, and I'll answer it from the perspective of a researcher, a publisher, an editor, and a clinician. It's like anything we do to really be sure we need evidence. We need research and evidence that signals to us as users and as a community that the agent is safe for independent practice. And then even when we have that signal, it's important to have ongoing monitoring, if you will, or surveillance. We just want to make sure that even after a particular agent is deployed, it's safe to do, it's accurate, and it has the best interest of the patient and of the end user in mind.

And to be clear, we don't have that kind of signal today. I think we're in this very strange situation where we're using tools at the bedside that are changing how we practice, but the really solid randomized clinical trials that show benefit and safety are still evolving, let alone regulation by an independent third party.

Dr. McDonough:

Jean-Claude, your thoughts?

Mr. Saghbini:

Yeah, I'm in complete agreement that we should let AI act in areas we feel confident that AI is safe to act in. I do want to pose, though, the problem that we're faced with, which is—and Brian, we talked about it last time, right—a hundred million Americans today don't have access to adequate primary care. Some don't have access to any primary care. And as such, there is a major access problem, and the only way we can get ourselves out of that access problem is through AI, because we cannot hire our way out of it. There are not enough

physicians and clinicians out there. And therefore, it is in our collective best interest as a country, or as humanity, to have AI be able to do more and more and act more and more. However, we need to make sure that we do that in the constraints of safety. So the urgency to what Raja just mentioned—the urgency to test, validate, and make sure that AI can act so that we can give AI more and more power to act—is extremely important.

Dr. Abdulnour:

One remarkable thing that got us both everyone excited, in particular Jean-Claude and I—full disclosure, we're friends—what got us both excited but at the same time cautious is the fact that these tools are able to solve problems that have plagued patients, and clinicians in particular, for a long time. The data shows that if you were to calculate how many hours a clinician must spend every day doing their work for a typical roster of patients in primary care, they need to spend 27 hours in a 24-hour period, right? So that's the classic study.

And so anything that can inch away at that, decrease the burden, and also just bring back the pleasure of taking care of patients is really something that will make a huge impact. And I think that's why people are very excited about using large language model-based clinical decision support tools—because they're doing something that's always been very hard. Knowing, what's the latest evidence? What's the latest guideline? How does that guideline help me today with the exact patient in front of me? Having a scribe and a phone that can draft a note for you, you don't need to jot down notes and look at the computer, all while the patient is there looking at you. These are low-hanging fruits that, overnight, became something that could be addressed by technology developers, including Jean-Claude and his company.

Dr. McDonough:

So you're bringing up tasks where AI can safely act. They're good examples. What are places where it should clearly remain advisory?

Mr. Saghbini: I'm going to put a lens on AI interacting directly with patients. We're able to deploy AI at scale to engage with patients, explain, check on patients, and do a whole bunch of tasks. It's not that it's replacing humans; it's that we've never had the human capacity to go do these things. So we're doing things that today aren't done and can't be done, and they are being done safely with AI in this patient engagement context. But we stop short of this diagnosis and treatment. And when engagements and conversations with lots of conversational AI start getting to that point, this is where a handoff to a qualified, licensed human is the right thing to do. And the smoother the handoff, the better.

I'll add one thing, which is historically, because we couldn't engage with patients at that scale, we actually did not know when patients needed to be interacted with by a human, right? So now, yes, AI is stopping short of diagnosing and treating, but we're actually able to identify areas of need that patients have, which historically had gone unnoticed. And when we find those, yes, AI does not treat them, but AI identifies them and actually does the right handoff and gets the patient down the proper care pathway. Raja, what are your thoughts?

Dr. Abdulnour:

Yeah, listen, it's a great question. There's a couple of things. The way I think about it is, it depends on the stake of the task, right? So if I as a clinician am using a particular tool, my level of reliance on that tool and my level of skepticism of the tool depends on the task. So for example, I'm in the intensive care unit, and I have several AI models at my disposal to support my decision-making—an AI chatbot, whether it's ChatGPT, Claude, OpenEvidence, you name it—any of these tools. I will never surrender my judgment if the task I'm asking the AI to help me with has a life-or-death implication on the patient. So if I need to change an antibiotic, if I need to decide on the next course of action on a patient in front of me, my level of skepticism is very high. And in the same way that in these high-stake situations, I will rarely delegate to a human clinician and will only want to get informed by them, I certainly will not delegate to an AI tool. And so I think it's really getting a sense of, what are the stakes of the task? So that's one.

And I think one good thing that technology companies are doing is sensing this and starting with, what are the low-stakes tasks but also high-impact tasks that we should address first? And Jean-Claude, I would love to hear—how do they approach this? And the second thing is, there's data that's showing that humans are very vulnerable to the fluency trap. The thing about these tools are that they sure sound like they know what they're talking about, right? They sound like experts. And when something sounds like an expert, you're much more likely to believe it than if it doesn't. And these tools, especially the great ones out there, are designed to please us, and they will never fight our instinct. And so when you're working with a tool like that, the likelihood that you'll defer your judgment to them is high.

And there's recent papers that show that if you instill an AI with subtle errors, humans won't detect it because, "Hey, it sounds an expert, so I believe it." So the challenge for clinicians is how not to fall for that trap. And I think there's some training that needs to be done for us clinicians on making sure we don't fall for that fluency trap.

But I'd love to hear from Jean-Claude. Is there something that could be done on the technology side to signal to clinicians and patients

that, "Hey, I may sound like an expert, but I'm not. Don't believe me fully. Make sure you ask your doctor first"?

Dr. McDonough:

So Jean-Claude, when you build an agent like Tom, have you worked things like this into it—those thoughts that AI wants to really please us?

Mr. Saghbini:

Yeah. If we think about Tom in this concept of primary care as a service, there are two areas that this agentic experience of Tom represents itself. One is patient-facing, one is clinician-facing.

On the clinician-facing side, it's always important to actually signal that this is a recommendation coming from an agentic framework and that AI is involved. Because physicians, or all clinicians, read information all the time in EHR, whatever it is. And the fact that this information was generated by AI versus human-curated and written is extremely important.

I think the other important thing is the reference to the information, because sometimes, it may be correct or may be incorrect, but regardless, where the information was derived from is extremely important such that you give the opportunity, when there is doubt, for that action to not be a multi-hour action of trying to go figure out where that is.

Now for this, we talk about workflow—from within the workflow, being able to very quickly get to the supporting evidence as to why a particular recommendation was made. The keyword is recommendation, right? Because we're still not making AI diagnose and treat versus clinical decision support for the clinician. So that's on the physician side. So again, always stating it's AI and always referencing it back.

On the patient side, we've been actually extremely successful at the guardrails. So when we think about our world, we think about it in guidelines and guardrails. Guidelines are what agents should be able to do within the scope of the task. Guardrails are the no cross lines. And an unbounded experience with any AI, you name it—with a large language model, it gives you this feeling the LLM wants to answer the question for you and cannot help itself but give you an answer. In the frameworks that we're building, we're very successful at putting the guardrails such that when you start hitting the line that should not be crossed, the LLMs are actually stopping and not crossing that line. The agents are stopping and not crossing that line. But then to test that at scale, actually. You have to use AI to test AI with millions of interactions to make sure that your agents are staying true to this and not crossing the guardrail.

Dr. McDonough:

So Raja-Elie, when a model gives a recommendation, what should a physician be able to inspect? The source, reasoning, the confidence level, the guideline basis, patient-specific data? Is it all of these things? What's your thought?

Dr. Abdulnour:

It depends on the question, but I wholly agree with what Jean-Claude is saying. Listen, in my mind, I've always approached trusting an agent and trusting a tool based on three things. One is the effectiveness or the accuracy of a particular tool. How good is that tool in doing something?

The second thing is transparency. Is the agent able to show their work? Are they transparent in how they got to where they are? And this is where the ability to cite or to point a user in the direction of the evidence and tell them, "This is why I'm making this recommendation, and this is where I got this information from," is essential. And frankly, it's one of the reasons why the *New England Journal of Medicine* has, as a group, decided that when the majority of US clinicians and millions of patients are using tools to address health-related questions, we need to make sure for our mission that these tools are grounded in the best information and that they can point to trusted information, which is why you would see NEJM and NEJM-related content in a lot of these tools.

And then the third piece is accountability. I'd love to know that the tool I'm using is great. I'd love to know that it can give me recommendations. But what I would really want to know is, does it have my best interest in mind, and it's not going to harm me? So do no harm and do good.

And an agent and a piece of software cannot by definition have an incentive. This is where the unusual and important collaboration between technology and healthcare providers and regulatory bodies and researchers is coming together because we have this agent that doesn't have any incentives per se. And so how do we define benevolence? How do we define accountability when we're using these tools? So for now, really making sure that clinicians are using tools that are grounded in trustworthy medical information, like NEJM and others, is really critical. Otherwise, it's going to be very hard to trace back and cross-check a particular answer.

Dr. McDonough:

For those just joining us, this is *The Convergence* on ReachMD. I'm Dr. Brian McDonough, and I'm speaking with Dr. Raja-Elie Abdulnour and Jean-Claude Saghbini about what makes clinical AI trustworthy enough for real-world integration, from medical evidence

and workflow design to governance, accountability, and patient safety.

And that does drive us into even a more detailed conversation. For instance, Jean-Claude, in the Tom episode we had, you described guardrails, and you brought them up again today around not diagnosing and not treating. But how do these guardrails actually work in practice thus far from your experience?

Mr. Saghbini:

So let me define what guardrails look like. Some guardrails are clinical guardrails. Some guardrails are personality guardrails of an agent that you create where you want the agent to behave in a certain way. And some are regulatory guardrails, right? They're not necessarily clinical, and so I'm going to talk about how they implement it, and maybe how they test it.

So the way we implement them is they're also written in the same language that we speak to LLMs with, except that we give them a higher order of priority. I'll give examples of those. A guardrail is, "You will not have a conversation without identifying that you are an AI agent." That's a very strict guardrail. We don't want that to impersonate a human. Another one is the one I gave earlier, which is, "A patient wants to change their medication. You don't prescribe medication. You don't change medication." So we write them with this higher order of importance, within the agent framework. And then what we do is, in order to validate that they work, we have to test them at scale through effectively brute force, trying to break them. For LLMs, by definition, their behavior is non-deterministic, right? You've had that experience where no question gets the same answer twice if you ask it. So this non-deterministic nature therefore needs to be tested through probabilistic mechanisms. And the way to do that is we have the power of computing.

So we then run at them extreme computation trying to break them, and then we react to where we are finding loopholes and then do it again until we get to, effectively, an agent whose guardrails are unbreakable.

Dr. McDonough:

I think of this all the time when I'm working in clinical setting, and it really comes down to, who should be accountable when an AI-supported workflow fails? Is it the physician? It was easy when it was just me working and making a decision. Boy, if I failed, I blamed myself. But now you've got health systems, vendors, model developers, or a combination. Where would the accountability lie? Is it still with the physician?

Dr. Abdulnour:

I'm not an expert in that area, but I've read enough and I've listened to some experts, and the key thing here is that, to this day, accountability lies in the hands of the clinician. The basics remain, which is to really make sure that harm doesn't happen to the patient. So that's critical. And then number two is to always keep in mind the standard of care. And, whatever tool we're using as clinicians, we need to do no harm, and we need to keep in mind the standard of care, which is why it still matters. We need to know the guidelines. We need to know the standards of care and the areas that we use.

AI is there to inform, make things a little bit easier, and give us more time to do other things. But at the end of the day, the decisions are ours. Brian, again, as a clinician, I take full responsibility for any decision that I make, whether it's informed by AI or a resident or even a consultant that I've called onto my team.

Dr. McDonough:

So with that in mind, I know with *New England Journal of Medicine*, for instance, you're doing a great job helping physicians make decisions and become aware. But how should a hospital or a practice leader evaluate an AI tool before deployment? Where do you go? What direction, and how do you learn? It's such a sharp learning curve. What do we do?

Dr. Abdulnour:

Again, the jury's out there, and the way I think about it is there are two arms. There's making sure the human is ready to use AI. What are the competencies of today's clinician in using AI? What is the baseline AI literacy that the clinician must have before using AI? Because as anyone who's used any of these tools, a single twist in the prompt can make or break an answer. Which tool to use and when, so on and so forth—there's a fair amount of AI literacy that's required.

And second, there is due diligence by the health systems when they're thinking of implementing an AI tool in a particular system. And that is, again, if I come back to accountability, transparency, and accuracy, the health systems need to look at all three. Is there evidence behind the efficacy or safety of a particular tool? Is there peer-reviewed research or the randomized clinical trial? Do the developers have model cards that can support how the data was trained? A tool could be absolutely fantastic, but it's been all trained from patients in Scandinavia and therefore does not apply for the rural population in Atlanta, for example. So there's that.

And then transparency, again, looking at the model card, is baked in the tool—some elements that provide transparency. Does the model show its reasoning when you ask it? Does it cite references? And for benevolence, if you think about the pharmaceutical industry

or even medical licensing, health systems benefit from third parties that can say, "This tool is FDA approved. This clinician is board certified." We don't have that in medical AI. And so that determination of accountability unfortunately is something that health systems and clinicians need to be on the lookout for sometimes when making their own judgment calls. But we, NEJM and others, are really thinking about this. There's a trust gap; how do we fill that trust gap in a world where the models are not only changing all the time, but just like Jean-Claude was saying, you can ask the same model the same question twice, and it's going to give you a different answer?

Dr. McDonough:

The title of this program is *The Convergence*, and I was really excited having the two of you together because we have a healthcare technology builder and a physician-editor working together. So when you are talking together—and you mentioned you're friends and you can collaborate—what should responsible collaboration look like? How do you see it? Because I really can't think of asking two better people than you.

Mr. Saghbini:

Yeah, it's a really good question. We actually have found that a new model of development is needed to create agents. It's not the old model of scrum teams and development teams, et cetera. So this is one of many, but I'll describe one I'm very familiar with, which is a squad model. And on that squad is a clinician with the expertise needed for what those agents are trying to do, whether those agents are in diabetes management or musculoskeletal or whatever that is, right? There are engineers on that squad. There's user experience on that squad. And effectively, it's a team sport. So the AI development now is a true team sport because you need to cover so much in what is ultimately just software.

At the end of the day, you're writing code, right? It doesn't matter how you write the code; you're still writing code. But that code needs to do clinical-like things, act like a human, and function like code—not break and scale and all of that—and be tested and go through regulatory approvals. So to be able to do that, you need that team approach. If I go back to your earlier question on, how does the health system think through it? It also mimics how a health system would do it—so to Raja's point, how do you evaluate it? How does a health system evaluate it? I'll add to what Raja said with this lens of three things. One is, does the AI do its job? Does it function? And you need to both validate the company that developed it as well as the tool. Does it accomplish the task? And each agent has a specific scope that it should be able to function within. So that's one lens of evaluating.

Second part of evaluating is, as we talked about, safety, right? So you need a governance lens. So we're seeing most health systems now have AI governance teams, and these teams are tasked with governing that AI is going to be deployed safely, can be checked, can be validated, and so on. And then the companies that are building it are going through these. These are the two lenses.

The third lens is not a technical lens. The third lens is change management lens. What would it take to input that AI in a workflow? Who needs to do what? Who needs to change their behavior? Who needs to be trained? The less retrained and the lower the bar in change of management, the more successful you're going to be.

So that's a macro view of a health system. In the micro view of the developers of agents, they need to have these things in mind. Can I make it function clinically in software, et cetera? Will it pass validation, regulatory, and all of these things? And will this agent be an acceptable agent in workflows that exist today? And what's the hurdle to get those agents to be accepted?

Dr. McDonough:

Great answer. Raja, I want you to jump in too, because what I'm getting from this is, we hear so much hype about AI, so we fear so many things—all the stories being covered and people are talking about. But what you're really talking about is the reality and breaking it down and figuring out how to use this logically. So Raja, what do you think?

Dr. Abdulnour:

Yeah, I love your reflection, on the title of the podcast, *Convergence*. And it does reflect when Jean-Claude and I started talking about this, and again, full disclosure, Jean-Claude and I have worked together for a little bit over a year. And thinking there's something generational here. For a long time, I considered myself a technologist. I love technology. You can see a bunch of tech-related nerd aficionados with gaming paraphernalia, but coding and software development have always been a little bit out of reach. And I've always appreciated AI and thought about it.

But there was something about these large language models. With these tools, I think the genius is actually not only in the model itself, but in the developers that built a wrapper around these apps that made them so available that all of a sudden, the skill gap narrowed. And so now you have these experts in their areas. And, am I the uber expert? No, but I do know a fair amount in my particular area. All of a sudden, I can use a tool, I can code, and I can do things that were very hard for me to do. But I can also see where it's good and where it's bad.

In early conversations with Jean-Claude, we came to the same conclusion. Jean-Claude is also a technologist who has a very deep and sound understanding of healthcare and healthcare delivery. And so that's the type of synergy that is happening all around us. And for us, a clear example is how the NEJM, a 200-plus-year-old institution, is now recognizing that we need to partner with technology for the benefit of clinicians and patients.

There's a new actor on the block here that is, whether we want it or not, an important new nexus of information—a co-clinician, co-scientist, or co-pilot that we need to work with. And so that convergence of skills—you're absolutely right. It gets me very excited because I love technology, I love medicine, I love science, and I'm mission-driven, and so it's a good moment to be there.

And one question that keeps coming up is, what's the future of clinicians? What's the future of skill and expertise? And I think one thing that's been clear to me is that expertise is more important than ever. The ability to know when a particular AI tool is going off the rails or being effective requires expertise, and there is no shortcut to experience and learning.

Dr. McDonough:

For a practicing physician listening tomorrow, let's say, what is one AI use case that is reasonable to explore now? What would be something they should, if they're getting on board, explore at this point? Either of you can answer that. Your thoughts?

Mr. Saghbini:

I think it's AI agents that engage with their patients to do the work that they don't have the capability to do or capacity to do. To Raja's early point on the 27 hours a day that the primary care physician needs to manage their panel of patients, using AI agents to engage with these patients—I think that would be the one I would go after.

Dr. Abdulnour:

A use case I've used—there are two that I've used in practice that I think are worth experimenting with. The first is a scribe. I don't know about other clinicians. I can tell you I don't write fast, and I don't type fast. And when I write notes, I can barely read my notes after the encounter. Having an AI scribe has been a terrific help. And this is one use case where data is evolving. We have some pretty good randomized clinical trials, which is very interesting because when you look at how many minutes you gain in terms of efficiency, it's barely a minute or two per encounter—not much. But despite that, the vast majority of clinicians, when you look at their wellness metrics—how happy they feel about using them—it's off the charts. And so clinicians are really happy about using them, even though it's not necessarily shaving off tons of time of their workflow. So there's something that these scribes are doing that's making clinicians happier. And another use case is clinical decision support, so asking about the patient in front of me, "What is the latest guideline about treatment?" And that has really changed how I do either.

And the caution on both—scribes will miss things. They will say things I didn't say. They will completely change the sentences. And they don't put in the chitchat. These tools, by default, by design, just don't include the chitchat between a clinician and a patient. And so if chitchat is important for you—chitchat meaning, "Hey, how's your neighbor? How's your family member? How's your cat? How's your dog?" The human aspect of the encounter is actually taken off of a lot of these, AI scribes. I can't emphasize enough the importance of double-checking the output. And then same thing with the clinical decision support tool. Sometimes it will say yes, but then the explanation is no. But for some reason, it's getting me much more engaged with the experience. I feel the need to double-check and to cross-check has actually made me learn more and get more engaged with the content. I can't put my finger on it yet, but we definitely need more research on that area.

Dr. McDonough:

One other follow-up here in this series of questions: what are questions every physician should ask before trusting a clinical AI tool?

Dr. Abdulnour:

There's a recent study in *JAMA Network Open* that looked at the determinants of trust in patients using AI. So this is not about clinicians, but at the end of the day, clinicians can be patients sometimes. So I think it was a good study that just highlights what matters to people. And what matters most is knowing that the particular tool is at least better than a human in doing a task. And so I think the first question is, do we have evidence? Do we have research—good, solid, objective research—that tells us the particular tool is effective and safe? And what's the research like?

These are important questions. Now, does any clinician at the bedside have the time or energy or know-how? No. And this is where the health systems come in. This is where a Chief Health AI Officer of a particular health system needs to weigh in. They even have now these AI competency committees that look at all the AI tools and give recommendations. That's one question up to the clinician to ask.

Dr. McDonough:

One final question for each of you because we've had a chance now to talk individually, together. Jean-Claude, I'll start with you. Is

there something I didn't ask you or you said, "Why didn't he ask this question?" that you want to bring up and talk about before, before we do wrap it up? Any final point you want to make for us?

Mr. Saghbini:

Yeah. What do the patients want? That's the hard part, right? Because this whole thing that we're in is all about them. So patients are consumers first. Number one, they're consumers. And as such, whether we give them tools or don't give them tools patients are going to go do things. They will go search for things, look for things, talk to their neighbor, or get medications with their cousin if they worked on their cousin. And so they are seekers of information about their own health.

We talked about the access problem that we have, so they're not getting enough care from the overall health system in the country. And they are starting to use AI. They will continue using AI. They are becoming experts in AI. The kids are AI native. So with that lens, what is their expectation? And I believe their expectation is to be given AI and for providers of AI to test it rapidly and give it to them. And they are trained on how to use it. They're using it for other non-healthcare things. They know when, if they ask for an architectural diagram to be changed, if it doesn't make sense or it makes sense. So they have intuition to AI, and I believe they are becoming demanders of AI.

The other thing is also what we know about humans—so patients, but humans—is that humans want whoever's giving them a service to be equipped with the latest and greatest of technologies to give them that service, right? In the quintessential example, I want the pilot to fly my plane to have as much technology as possible at their disposal. So although I probably don't want AI to make a decision on behalf of my physician, I would be worried if my physician was not using AI. And that's an interesting line. It's not a gray line. It's actually a very direct line. It's, are they using AI or is AI doing the work? We can keep aside that AI is diagnosing and prescribing, but to the extent that physicians are not using AI, for me as a patient, I would feel I'm not getting as good care as I could be getting.

I'll give one example. Raja Elie, I'd love to get your thoughts on the question earlier, but I'll give one example. Two weeks ago, I went to a physician for a shoulder problem, and it was the first time I had an amazing experience. It was my first time with a full scribe; she did not touch the EHR. She walked into the room and put her phone down and then left the room, and we spoke 100 percent of the time. Now, knowing about scribes, I knew what I was listening for a bit and that she was saying certain words for an AI scribe to record them, et cetera. But it was the most amazing experience. I'm like, "Okay, that's what this technology is doing. It's actually getting me closer to my physician than further away."

Dr. McDonough:

I've got to tell you, I almost find myself now, because I'm doing this, looking at the patient too much. Does the patient think I'm weird because I'm staring at them? But it's great because you pick up things. My example is a mother with two children. I was examining the child, but I could tell Mom was overwhelmed, and it wasn't even mom's visit.

And I was like, "Are you okay?" But when your nose is in a computer, you don't see Mom. You just are saying, "Okay, height, weight, blood pressure?" Then all of a sudden, I said, "Are you okay?" "No, I'm really not." And everything shifted in family practice to the mom. And that was probably more important for the kids than the kid exam.

So I think you're right on. And I think patients are going to love that because I hear complaints: "I never talk to my doctor anymore. I feel like I'm just a number." And I think it is the eye contact, which you were talking about, that matters.

So let me ask you, Raja. What didn't I ask you? What do you wish I had brought up?

Dr. Abdulnour:

So I wonder, Jean-Claude, if you have the same doctor, because I had, a few months ago, an injury in my wrist. So I went to see an orthopedic surgeon, and she came in, put down her phone, and then talked to me. And I was like, "Wow." I was blown away. And that's what struck me. Then she started telling me how it changed her practice. It gave her more time to take care of her kids, it gave her more time to do research, and so on and so forth.

One thing that's really interesting, Brian, is knowing that there's a scribe, and—this is where I put my educator hat on—for the scribe to work well, you need to think out loud. And as you think out loud so that the scribe captures your reasoning, two things are happening. One is you're actually deliberately reasoning out loud, and so from a purely deliberate practice, think about it as a student, resident, or trainee—they now need to think out loud for the scribe to work well. But it's also offering to the patients a window on how we're thinking instead of waiting to see the note coming at the end. And then this is where maybe patients are like, "Why are they thinking out loud this way?" I think it's a fascinating world, and from an education lens, it's really critical.

One thing you didn't ask me is what keeps me up at night. And one is the health of my kids; I just want to make sure my kids stay healthy. The other thing is, if we thought that misinformation was a big problem a few years ago, the risk of misinformation right now is exponential. Think about it. Up until AI, humans, and in particular journals and editors, were the ones saying what is important in science

and what information is relevant. Editors were writing about it, and then editors were publishing about it, Brian, what you're doing in your podcast, you're selecting the topics, you're having the conversations, and you're disseminating it. But then social media comes in, and you have these algorithms that are deciding what people should watch and should not watch. When you pull up your phone, it's not a well-meaning editor that's telling you what to look at. And so that step has become taken by an agent, and now you have tools that can actually write the content for you and even make it do videos for you.

So what keeps me up at night is that the health of our information ecosystem more broadly, but including health information is at risk of being taken over completely by software and by code, whose incentives are very different than mine—mine being taking care of patients and making sure the physicians do their job well. And that's what keeps me up at night, and that's what wants me to make sure that there's always an editor in the loop, right? We make sure that they all have our content so they can ground their information or content. We solicit the best research out there on these tools. We play a role in regulation. We play a role in our voice. So that's frankly what keeps me up at night: the promise and the peril.

Dr. McDonough:

Wow. That's a great place to leave it. I want to thank my guests, Dr. Raja-Elie Abdunour and Jean-Claude Saghbini, for joining me to help clarify what trust should mean as clinical AI moves from answering questions to participating in care workflows.

Raja-Elie, Jean-Claude, it was fascinating having both of you back on the program. This was a lot of fun, and we really talked about some critical issues. Thanks to both of you.

Dr. Abdunour:

Thank you, Brian.

Mr. Saghbini:

Thanks, Brian.

Dr. McDonough:

For ReachMD, I'm Dr. Brian McDonough. To access this and other episodes in our series, visit *The Convergence* on ReachMD.com, where you can Be Part of the Knowledge. Thank you for listening.