

Take It or Leave It: Episode 50

(This transcript was generated through AI technology.)

Josh Seidman: Hi, everyone, and welcome back for the 50th episode. That is right, the 50th episode of Take It or Leave It, where we discuss the hottest topics in the world of workplace leaves, absence management, and accommodations. I'm your host, Josh Seidman.

What a time it is to be a soccer fan. The World Cup is in full swing—or kick, I suppose—which is a dad joke if I've ever heard one before. And I don't know about you, but I have been glued to every match that I can catch. This tournament has already delivered drama, heartbreak, controversy, elite goal scoring, and so much more.

Let's start with some shockers. Germany and the Netherlands, two of my favorite countries and two of the traditional giants in the sport, were both bounced in the round of 32, and both on penalty kicks, no less. I mean, talk about gut-wrenching ways to exit a tournament. On the flip side, France has looked absolutely dominant so far in the tournament—just so composed and firing on all cylinders. And of course, there's been some individual brilliance as well. Messi is still doing Messi things. Mbappé is looking every bit like the best player on the planet. Haaland has been finding the back of the net game after game. And Kane seemingly single-handedly kept England in the tournament with his expertise in their round of 32 match. As I said a minute ago, the goal scoring at the top of this tournament has been a real treat for us fans.

All right, enough of my soccer soapbox. Let's get into today's episode. For today's episode, we are continuing our ongoing exploration of one of the most important and fastest-moving topics in our leave and accommodation space. That is artificial intelligence and its growing role in leaves of absence and accommodations. Listeners might recall episodes 43 and 44, where I had the pleasure of speaking with Brian Bass from the Disability Management Employer Coalition about ethics, guardrails, and their think tank work happening in our industry. Today, we are picking up that thread and bringing in a technology and industry vendor perspective to complement those earlier conversations.

To continue this discussion, I am thrilled to welcome to the podcast not one, but two fantastic guests from FINEOS—Patricia "Trish" Zuniga and Lori Welty. Trish and Lori both recently authored a really thoughtful blog series titled "Man vs. Machine," which digs into how AI is being used today in the leave and accommodation space, what it can reasonably do, and where human involvement and discretion still need to live. I also had the good fortune of catching a live presentation that they gave on this topic at a recent DMEC Compliance Conference event, and I knew right away that I had to get them on the show to dive in deeper.

Trish is an Integrated Disability and Absence Management—IDAM—Compliance Manager at FINEOS, dedicated to bridging the gap between complex legal mandates and innovative product compliance. She provides critical compliance support for FINEOS's IDAM product, serving as a subject matter expert in FMLA, USERRA, and state PFML regulations. Trish also plays an active role in the company's AI governance initiatives, helping identify emerging regulatory risks and supporting the responsible use of artificial intelligence. She has served as a speaker on PFML, AI, and related compliance topics for the Society for Human Resource Management (SHRM), the DMEC, state agencies, and broker forums. Trish holds law degrees from both the George Washington University Law School and Ateneo de Manila University School of Law. She further holds the Certified Leave Management Specialist designation from DMEC.

Lori is a nationally recognized leader in leave, disability, and workplace accommodations compliance. As SVP of Compliance at FINEOS, she leads strategic product compliance initiatives, partnering to shape product strategy and develop scalable solutions that help customers navigate an increasingly complex regulatory landscape. She advises on emerging federal and state leave laws, PFML programs, the Americans with Disabilities Act, and other evolving workplace compliance requirements. Lori specializes in bridging the gap between legal requirements and real-world implementation, helping organizations build compliant, effective solutions. Lori earned her JD from the University of Colorado Law School. She received her bachelor's degree from Johns Hopkins University.

Trish and Lori, welcome and thank you so much for joining us today on Take It or Leave It.

Lori Welty: Super excited to be here.

Trish Zuniga: I'm so thrilled to be with you both.

Josh Seidman: As I mentioned in what felt like an ongoing would-never-end intro—sorry about that. I really like the World Cup and you both are so impressive that I had a lot to say at the get-go there. But I'm so thrilled. I've been really wanting to have this conversation with you both because I heard you speak a few months ago at the DMEC Compliance Conference for this year. So I'm really happy that we're making this happen today.

Lori Welty: Great, great. Well, thank you.

Josh Seidman: So lots of ground to cover. I'm ready to hit the ground running. So Trish and Lori, a couple of background intro questions for you. Before we dive into the meat of today's conversation on AI, I would love to just give our listeners a little bit of background to get to know you both better. Can you each just share a few thoughts about your professional backgrounds and specific roles at FINEOS, and then also just tell our folks a little bit about FINEOS as a company and the role it plays in the leaves of absence and accommodations space, including how you all fit into the broader leave and accommodations ecosystem.

Lori Welty: I'll kick that off, Josh. Yeah, thanks. My first part of my career was kind of a traditional legal path. I started the first decade of my law practice as a corporate and employment litigator, so in the court setting—so did that for the first 10 years or so. But over the past 15 years, I've been focused exactly on the absence and disability space. So during that whole time, I've obviously been in a compliance role as an attorney, but really focusing on that intersection of law and technology. So it's an interesting role that I have really learned to love.

So my role at FINEOS is ensuring that our leave and disability solutions help organizations administer their programs exactly as the law requires. So, you know, at FINEOS, I lead the product compliance, making sure that our software implements federal, state, municipal rules—from PFML to ADA to FMLA—so that our customers can manage these programs accurately and confidently at scale.

And so just generally about FINEOS—probably should have started that before I went into what I do there—but FINEOS is a global provider of absence, disability, and claims management software. So we are a platform that is designed to help insurance carriers, employers, third-party administrators who need to administer leave, disability, and accommodations at scale. So, you know, in terms of the ecosystem that you mentioned, I view FINEOS as really sitting in the middle of that ecosystem. So we sit in between—in the intersection of employers, carriers,

TPAs, and employees—right in the middle of that, providing tools that allow those requests to be handled compliantly and, in light of today's conversation, with the right blend of automation and human oversight. So our platform does everything from intake, eligibility, time, payment calculations, case management, communication, and of course, overlaying all of that is the regulatory compliance that is part and parcel of all of that.

Josh Seidman: Wonderful, thanks, Lori. And Trish, how about you?

Trish Zuniga: Sure, I'm Trish Zuniga, Compliance Manager at FINEOS. So I entered the leave and absence management field in 2016. This was exactly when paid family and medical leave was expanding across the US. And so I like to say that I'm PFML native, just as I am AI native. So I've seen the explosion of all of these new and amended leave laws that transformed what leave compliance looks like. And it's such a dynamic area of labor and workforce regulation.

And how I fell into this field is kind of backwards because I had a background working in the legislative and executive branches of the Philippine government. So it gave me this really strong foundation to see and read how laws are developed and implemented and enforced. So applying those skills—that policy and regulatory perspective—that's how I help companies navigate what the laws require today and also to anticipate where the next regulatory trend is coming next and help prepare for that change instead of just reacting to it.

And what I do at FINEOS, working with Lori, is we take those complex rules of leave and administration and we turn them into something that our software can use. Now, we're not technical people. We're legal and compliance people. So what we do is we translate those legal and compliance requirements into those triggers, conditions, and workflows, and building it into our software so that the system can do things like automatically check eligibility, calculate benefit durations, and track intermittent leave. And so what we're sharing today comes from working directly in the space where these configuration decisions actually get made.

Josh Seidman: That's great, Trish. Thank you both. Really interesting backgrounds and really helpful explanations on both of your roles within the company. I love some of the comments that are leading us into the rest of today's discussion—you know, the blend, the right blend of automation and AI, translating compliance requirements into the software world. I think those really set the table for our next set of questions here.

I mentioned this in my intro. You've done a really thoughtful blog and some presentations in recent months called "Man vs. Machine." And you were examining, again, how AI is used today in our leave and accommodations bubble, what it can do, where human involvement and discretion still need to live, and among other tentacles and concepts. As I mentioned, I was also really lucky to see you present on this topic live a few months ago. So I want to start by actually taking a step back from that body of work. What are some of the overall takeaways or trends that you are seeing right now in terms of how AI is popping up in our industry?

Lori Welty: Josh, one of the things that we are seeing is that organizations are really looking to experiment in lots of different ways. They're trying to understand where AI can make an impact, where AI is creating some efficiencies for their employees. In other cases, where it can automate some decisions—which we'll get into whether or not that's such a good thing—you know, where they can add in some chatbots and maybe speed up some people getting answers to questions that they kind of see over and over and over again, to speed up those answers. So we're seeing a lot of experimentation, just trying to understand and implement ideas to see: does this work? Does this not work? Is this helping us?

We're also seeing, for better or for worse, some rogue implementation of AI, you know, where all of us are in the AI ecosystem. Whether or not we work anywhere, there are students, retired people—everybody lives in the AI ecosystem now, so it's infiltrating every part of our lives. So we're seeing AI in a formal context, of course, actually being implemented in the workplace and required in some cases, but we're also seeing it pop up in really unexpected areas where people are using it on an ad hoc basis that maybe don't even realize that the way that they're using it is a concern at all.

So on the one hand, we have all of that kind of innovation and experimentation. And on the other hand, what we see—and we think this is healthy—is we see leaders and of course compliance folks becoming more aware of the risks of this implementation—concerns about privacy, transparency, human oversight, bias—and a little bit of a pullback to say, okay, this is exciting, this is new, this is interesting, but let's take a breather before we go all in on this.

So, you know, I think on a global level, what we're seeing is an acknowledgement that AI is here and it's a reality and there's nothing we can do about that. And I don't think anybody—well, I don't think most people—are looking to prevent that, but the acceleration has been so fast that by the time we have our handle on it, it's accelerated again. So I think there has been a growing focus, and we think this is a healthy focus, on responsible adoption and really understanding where we need to put frameworks in place so that we can allow for that acceleration and mushrooming, but have that framework in the background so it can happen in a responsible manner.

Trish Zuniga: That's a great point, Lori. And I think the only thing I can add to that really is to say that when we put out our blog series and when we came up with our presentation, it was titled "Man vs. Machine." And I think the trend now is really to see that there isn't a dichotomy—that man is competing versus machine. You know, skipping ahead to the end of our presentation—surprise—the theme is really man with machine, and how can human work be accelerated by the use of artificial intelligence, but in a responsible and thoughtful way.

Josh Seidman: I love a good plot twist. That was great, Trish. I think the background that you both just hit upon is what we all are experiencing—to a point that Lori made, everyone is in this AI ecosystem together in many facets of our lives and figuring out how to hold on and try to wrap your arms around it, but to do it in a concerted way, because of how fast the acceleration is, is important whether you're thinking about it in your professional life or your personal life. And that was why the article series that you both wrote and your team wrote stood out to me, because that was some of the themes of those writings. So thank you both for setting the stage for everybody.

I want to jump to a question from part one of your series, where you talked about this concept of AI fluency for leave and accommodation professionals. What does AI fluency actually mean in practical terms for the folks doing this work day-to-day? And why is building that fluency so important today instead of a few years down the road?

Trish Zuniga: So it might sound a little scary when people hear the term AI fluency. They might think, oh, this is such a lofty goal to shoot for. They might think it means learning how to code or becoming an AI expert. But really, what it just means is the ability to work effectively with AI and not just use AI. So it's not just about plugging in questions into an AI platform or knowing what the AI platforms are and how to create prompts. It's really—can you apply the AI to your workflows? How do you improve outcomes with AI? But also, know when AI is wrong or risky.

So it's important to know right now because AI is just something that you're expected to be able to use in your work. So I like to say that it's the same as using spreadsheets. You know, decades ago, you didn't have to know what an Excel spreadsheet was. And now you're just kind of expected to know how it works, how to use formulas, when the formula doesn't work, how do you fix it? Where did you go wrong? How was it calculated wrong or in a manner that doesn't get you the accurate information that you need?

So in the same way, if you're an AI-fluent professional, you know when AI can help you streamline your work or automate your repetitive tasks, summarize information. It means having to use your critical thinking and knowing how to validate your outputs, recognize those risks related to accuracy, privacy, compliance, and governance. Again, it's important right now because AI tools are becoming part of just everyday work across nearly every function. We're seeing job descriptions, job hiring ads that call for AI fluency, even for non-technical roles. So it doesn't mean that only AI engineers need to be AI fluent. Somebody in sales or finance or even a legal and compliance professional needs to be AI fluent.

Lori Welty: You know what that reminds me of, Trish? I have college-age kids, and when I've been interested in kind of seeing what they're learning in school, I've seen an interesting transition over the last few years about how AI is expected to be used or not used in their classrooms. And it seemed like for the first several years of their college experience—and in some classes I think it still is this way—there is a rule of no AI allowed, period. And they have to attest that they did not use AI, and their work is run through AI detectors, and there is just simply a zero-tolerance policy for AI, which made sense for a period of time. And I did understand that, but in some classes, that's still the rule.

But what I've been interested to observe is that they have some classes that they take now where they don't have to use AI, but it is expected that many of them will use AI. And so instead of having a rule that you can't use AI, you have to explain specifically how you used AI, how it impacted your work product, and provide that information. And so I view that—you know, like what you're talking about is the idea of becoming AI fluent, understanding where it is acceptable in our new society that we're in right now where AI is infiltrating things, where it is acceptable to use AI. You know, did you use AI to—did you type in a prompt and ask for a brainstorming session? And then did you use the brainstorming session from there and write your paper? Or is there some other way in which you use it? So I think there is kind of a growing recognition that AI has a place in a lot of different contexts, but it's so important to understand how you're using it, not just that you're using it. So I think that's a big part of that AI fluency.

Josh Seidman: Really insightful responses from you both. You know, it's interesting. Part of my background, which I've spoken about here and there on this podcast and otherwise—you know, I was a journalism major. I enjoyed reporting when I had that opportunity. I also studied physics and worked in a lab for a bit, and I feel like both of those disciplines and the way that you go about the scientific method and testing out different hypotheses and theories and trial and error, and then from a journalistic perspective, understanding what is a trustworthy source and digging deeper—I think those concepts feel very much like what AI fluency needs to become for a lot of folks. That you need to not only just operate at the surface level, but you need to think critically and take a closer look and go through that process of, if it didn't get the right response initially, try again, rework the prompt, ask AI what you can do to maybe solve for the potential issues that you're encountering.

So really interesting feedback, Lori, about what you've been observing from your children. And Trish, I think your explanation of AI fluency was spot on. So thank you both for that. That was great.

One theme that comes through pretty strongly in your writing—and we've already touched upon it quite a bit in some of these initial conversations today already—is this healthy caution around relying too much, too heavily on AI. From the employer perspective, let's start there. What are some of the biggest concerns that you would flag when an organization is leaning too hard on AI tools in the leave and accommodation process? And I'm also curious if you can flip that 180 and think about, what should employees be aware of when AI is part of the decision-making chain that impacts their leave or accommodation?

Trish Zuniga: As a compliance professional, I think we're very risk averse, just very naturally. So when we—

Josh Seidman: My wife doesn't love that about me, Trish, in certain settings.

Trish Zuniga: But it's necessary because on the one hand, using AI, it promises efficiency, automation of repetitive tasks, freeing up the human to do the more valuable work. So that is a great side of AI. On the other side—and that's the side that we live in and that we have a responsibility to point out—there's risks in using AI. There's three major risks that I like to call out.

First, AI operates at scale. So if something happens, if there's an issue in the model or logic, it gets replicated consistently and rapidly across the population. It's not this one-off legal case that you're able to deal with just the one time. No, it happens across all of the model.

Second, AI reflects the assumptions embedded in its design and data. So garbage in, garbage out. If the data that's being fed into it is outdated or not really accurate—there's gaps or uneven representation—those patterns may be learned, reinforced, and operationalized.

And third, in this leave and accommodation space—very important—it doesn't inherently account for those edge cases and non-standard needs. The problem is that AI uses kind of the majority rule. So what's the general pattern? What is efficient? What is the majority outcome? The problem is that leave and accommodation lies in that exact space where edge cases or non-standard needs are. So that's where we are. So using AI is kind of a tightrope between balancing risk, but also efficiency. And so the way that you operationalize AI is you put responsible AI into practice. And that's what Lori and I keep talking about in our day-to-day work and in presentations across the industry.

Lori Welty: Two things come to mind, Trish. And I was just thinking, Josh, you're mentioning about your prior work study in physics and scientific method. One thing that I think is so important when we think about AI and the way that we're using it is that we continue to use all of the same thought processes and kind of critical thinking that we've always used, and that we don't—when AI is in the mix—that we don't just throw them out the window and think that we don't need them. So I think the same level of curiosity and interrogation that we would apply to anything else should continue to apply when we're using AI. And I think it's just, you know, part of that fluency and part of a training that goes on in an organization to help employees remember that AI isn't magic. You know, it's a tool, but it's not always right. And it has to be looked at critically with our human brains as well.

And I think—you know, maybe I won't call anybody out in my family, but we all know somebody who has an over-reliance on AI when it comes to maybe asking for directions or coming up with ideas of things to do on a weekend. Actually, I will call out one person, which is a friend of mine who took me on a hike recently using AI, and she swore it was going to be a loop. And she swore that we were gonna end up back at the start. And we absolutely didn't. And about three hours into this hike, I said, I'm looking at the map. The whole time she said, look at the map, I swear. I looked at the map, I was like, this is not a loop. And so it ended up we had to do it three miles back. And that's just kind of an example of AI is not foolproof. Like, it is amazing, and we're so excited about all the things that AI can do, but there's nothing wrong with checking the map every now and then too, and just saying, let me look at the map. What's the harm in that, to make sure and double-check our work too. So I think just all the same challenges that we normally would use in our life to make sure that things are on track and working correctly, we should not throw them out the window just because we're using AI.

Josh Seidman: I agree. Really insightful, again, from you both. The hiking map story, Lori, it reminds me—this is years and years ago—one of my best friends and I, when we were undergraduate days, did a trip to Europe, and he refused to allow us to use maps on the trip. His saying was always, "Forget the map, I'm the map." And I think it was like a harken back to like The Mummy movie, a Brendan Fraser quote, I think, from even years before that. And I viewed that with skepticism then. So we did tend to get lost here and there. So sorry that you got lost, but glad you found your way back home on that recently.

Lori Welty: It was a fun day anyway, but—yeah, yes, exactly.

Josh Seidman: My next question. So using AI as a resource, or maybe as a shortcut depending on your perspective, certainly can be tempting and understandably so. You know, let me just ask ChatGPT or one of your company's approved AI tools what the FMLA rules are on this topic, or let me have AI draft this denial letter. But from both of your perspectives—and the last response got to this quite a bit—where is that line between using AI to work smarter versus using it to cut corners in ways that could actually get an employer or a claims professional into trouble?

Lori Welty: Yeah, I think here, one of the really important things in this space in particular—leave and accommodation space, which is so highly regulated—is to think about how close whatever the output of this AI request is, how close it gets to influencing an employment decision or an employee's legal rights. And the closer it gets to impacting a decision or a legal right, the greater the need is for human oversight, interrogating results, testing them, documenting them, and having that governance that we've been talking about.

Trish mentioned some of the efficiency, reducing administrative work. I think AI is incredibly helpful for things in the efficiency space—summarizing documents, organizing large pieces of information. Even, I think, in drafting communications, if we're talking about a first draft that's then going to be reviewed by a human, I think all of those things are really great uses for AI. But I think that line could be crossed when the AI is substituting for a legal or factual analysis or decision-making. I think it could be looked at as a tool in that analysis. And certainly, if you're looking at a large body of documentation and asking the AI to help you to assimilate that and understand it, it is part of your factual analysis, but it in and of itself cannot constitute the beginning and end of your analysis. It's a tool in there. So that's one thing.

I also think that at the end of the day, when we think about the fact that AI is impacting a decision, when that decision is made, if somebody asks that legal professional or leave

professional, why was this decision made? Why is this the outcome? They need to be able to explain that very clearly. The answer can never be, "we put it into this machine and out popped the answer." That's never gonna be an acceptable answer. So yes, you may have used AI to get to that answer, but you should be able to explain it without that. You should be able to say, you worked for this amount of days, and this is the law that applies, and here's how that math comes out. And yes, there may be something behind that that helped you get to that answer, but you should be able to look at the facts, look at the law, and explain that really clearly. So if at the end of the day, you're struggling to explain how your answer came to be, then I think you probably crossed a line in the way that you're using your AI.

And I feel like if you just rely on "the AI made me do it," or even "I'm not responsible for this, it was supposed to be the vendor or the IT person that checked it"—like, that's something that's going to get you into trouble immediately. And I feel like, Josh, when you hear this defense in court from the other side, you're going to be like, aha, my time has come. I'm winning this case for sure.

Josh Seidman: Yes, I hope for any opposition's sake that they have done their due diligence, but I hear what you both are saying. And it's, I think, a spot-on analysis and really practical to have the user reflect on: is this giving a legal decision? Is this impacting an employee right? How would I respond to the individual if they are questioning the outcome? Using those as some hurdles to go over and sort of stopgaps in the process. So really, really good segment there, both of you. Thank you for that.

I think, Trish, one of your responses a little bit ago mentioned in passing the concept of something that is somewhat underappreciated but really should not be in this whole overarching AI and the leave and accommodation space, or even more broadly, which is the art and science of creating prompts. There is a real risk of asking leading questions, getting back exactly what you wanted to hear. I think you made a comment before—you get in, you get back what you put in—which is not the same thing as getting inaccurate or a complete answer. So can you both just talk for a few minutes about the trial and error involved in creating AI prompts and how someone using these tools can push back on themselves effectively?

Lori Welty: Yeah, there's a few different things that I think about with prompt creation. First of all, I love this topic. I think it's such an important topic, and I don't think I'm breaking ground in saying I think this should be a topic in and of itself in terms of courses that you can take on it. And I think there are courses you can take on prompt creation. I think it's so critical to what we get out of it.

So I think of prompt creation, first of all, as very much iterative. So I think that we can start with what we think might be a good prompt, and then when we get the answer, we can rephrase that. And I know we've all had the experience of typing in a question to a chatbot or whatever, an AI, coming up with an answer, and you look at the answer and you think, I know that's not the answer. And then you say to AI, "well, I'm aware of this, and I know this is the case, so how would you respond?" And it'll come back and say something like, "you got me. Good point. Let me think that through again." And it'll come out to something else. So we know that there are ways of rephrasing questions to get at that. And so first of all, I think really the recognition that prompt creation is iterative—our first prompt is not going to be our last one. That's the first thing.

I also think—you mentioned the art and science of prompt creation—and I really do think it is an art. I think there's a lot of creativity that goes into it. And I think that there can be a lot of kind of individual thought processes. We can all have our own art to this.

One of the things that—before I'm done here—one is I talked about at the presentation that you were at, at the DMEC presentation, the concept of sycophancy in AI responses. And I think it is so important that everybody learn about AI sycophancy and read about it. I think it's just a training tool for people to understand what AI sycophancy is, why it exists—I mean, there's real reasons why it exists—and how it impacts us. And not just the fact that the way that we ask AI questions is going to impact the responses—it wants to please us, as crazy as that sounds, because we know it's not human, but it does want to please us. So that's the first thing—is that we need to be very critical of that.

But the second thing is that humans perceive flattering responses; it makes them more likely to believe those responses. And even when a human is aware of AI sycophancy, it actually makes them less likely to admit that they're wrong, even when they're educated on this topic. And so there's been studies done on this. I think it's really important to understand and read more about that and understand how AI sycophancy is impacting us.

But along that point, I think some of the very practical things that we can do when we are creating AI prompts are things like asking the AI for counterpoints. Tell it that you want it to argue with itself. Tell it that you want the sources. Tell it, you know, if you're asking for any kind of conclusion, say, "provide citations for every substantive statement that you provide." Ask it to provide neutral framing. Say, "I want to hear the risks and the benefits. I want to hear the other side." Also ask it to identify the uncertainties. You know, when you're providing this answer, provide the areas that you aren't sure about. Call those out for me. So I think really demanding of the AI to identify its weaknesses and the counterpoints can really help us to also avoid that AI sycophancy and avoid us from kind of feeding into what it's telling us too much and reading in what we want to hear there. So I think those are important.

And then one last thing, I think in terms of what we're asking for the AI, I think asking for the AI information brainstorms, bullet points, instead of asking the AI to produce a fully formed artifact—I think if we can keep ourselves really firmly in that space of creating the artifacts, but we're using the AI to help brainstorm and create those intermediary steps in the process, can help us to keep our brains engaged and getting a full picture of really what we want as our output.

Trish Zuniga: Yeah, great points, Lori. I think the only thing I'd add to that is, going back to the beginning, even before you put in your prompts, what I like to do is, just as a compliance professional, if I'm trying to summarize a law or pull out a specific legal provision, ask what it means—I like to take that information, I dump in the law, and I say, "don't do anything at this point, just read it." So I think it's the concept of retrieval-augmented generation that I also want to impart to the audience. The AI is going to tell you generalized answers. So if you ask for PFML provisions, it might pull in from summaries. You don't know where it's pulling from, actually, unless you tell it to provide its sources.

I like to make sure that the data that it's pulling from is something that I trust already. So I'll dump in the law, the regulations, the specific documents that I need referenced. And—we'll talk about that in a sec—make sure that it's a secure AI model. Don't just dump in the information on a consumer-grade platform, but make sure that there are guardrails to the data that it's pulling from. And then that's when I start iterating my prompts and pulling information and bouncing ideas off of the AI.

Josh Seidman: I have nothing else to add about both of those. Insightful, thoughtful, practical. So thank you both. The next question is related in terms of using AI in a practical sense,

depending on who the stakeholder is, right? Because the right use cases can look different depending on where you're sitting, what perspective you're coming at the issue and the topic from. From an employer or HR perspective, what uses of AI genuinely make sense today? And then let's also think about this question then from the employee's vantage point and where using AI could be helpful or empowering. And then also, there's another relevant group of players in our space, which are very importantly the carriers and the TPAs—the folks who are administering these claims at scale. Where for them do you see the most promising, responsible use cases?

Trish Zuniga: So just from a general perspective, AI is helpful or empowering when it automates genuinely repetitive tasks, when you don't have to think about it, when it doesn't make any consequential decisions. Where it does make consequential decisions, it still makes sense to use it and it's still useful, but you just have to remember that guardrails have to be in place so that it doesn't overstep where human judgment really should be making the decision and not the AI.

So for example, in HR, HR processes touch a lot of employee decisions—so hiring, firing, discipline, all consequential decisions. There are AI use cases at every phase of it. So from talent acquisition and recruiting, there's tools for candidate sourcing, screening, and matching, analyzing resumes. Performance management, where you can have someone's performance being analyzed across multiple data sources. In the leave and accommodation space, there's document intelligence for reviewing medical information, and rules and recommendation engines for eligibility and accommodation options, and case management tools. So all of these are very useful AI use cases. Again, the trick is to have guardrails in place so that they don't go sideways on you.

Josh Seidman: Thanks for that, Trish. Quickly, on the flip side of those use cases—and the tail end of your last response started talking about this a bit—what would you say are some of the biggest risks and pitfalls that you would call out for each of those groups that I mentioned: the employer, HR, claims administration folks? And I'm particularly interested in the ways that well-intentioned uses of AI can quietly go sideways.

Lori Welty: Yeah, absolutely. So let me start with the employees. So one of the things—and Josh, you may remember this from the presentation that Trish and I did back at DMEC—is the way that we had structured our presentations, we had a whole bunch of questions that employees or employers, HR, might be asking to AI. So just kind of common questions that you might ask, the way that you might use AI, and see what kind of answers are generated and how that might differ depending on the perspective.

And one of the reasons why we came up with this presentation is in talking to our clients about the way in which they are being approached with AI answers from their clients' employees. So, you know, we see employees who know that they're going to need a leave of absence and they go to AI and they say to their AI, "I'm having a baby in November and what are my leave rights?" And the AI, of course, will always come up with an answer. And so I think there's a million different ways that that can go wrong. And preparing for our presentation and doing our presentation, we explored some of those ways. But there's a lot of different ways that AI can give incorrect answers. And of course, one of the things that's so tricky about that is AI is so definitive. It's so authoritative, so easy to believe it.

And so I think the first thing is, if you're an employee, as we talked earlier, just really using your critical thinking skills and using it as a starting point to understand what your rights might be, but

using your traditional sources as well. So let it get you started on the right road, but don't believe it entirely. Continue to use the sources that you should be using and also reach out to your HR and your employer to understand that from their perspective.

But then from their perspective—same thing on that side—is making sure that when they're asking those questions, that they're using it as a starting point. So in terms of that initial part, doing that research, using it as a starting point, making sure that you're doing a deeper dive. But then when it comes to things like anything that can involve an employee's legal rights and employer's obligations, that you are always keeping that human in the loop. It is so critical there to be having human oversight and ensuring that the answers that you're getting are really appropriate for the exact circumstance that you're working with, because it's not so much always that the answer is wrong. It's just that you haven't provided all of the right information. So making sure that you have the full, full context there.

The other thing that I think is interesting—you know, a way that I've seen people use AI that puzzles me a little bit is for things along the lines of math. Math really doesn't need AI. Like, we have calculations that are so specific. So certain things like how's the benefit calculated or how much time do I have left? I'll be interested to see if we get reaction on your podcast in this—I may be an exception here, minority—I don't think we need AI to help us with those types of things. We actually have calculation engines, we have eligibility engines, we have entitlement engines, and they are designed—very carefully designed—with all kinds of input from the very specific information from the state regulators, and they take so much time and effort to go into these calculators. And once we have them made, let's use them. So I don't think we need AI in those areas. And I think it's much more likely that we're going to get wrong information when we use AI for things that are, as I want to say, simple, but as regulated as math in those situations. So I don't think AI has a place in every single aspect of this process.

Josh Seidman: Really, really helpful. Math is a universal language. There are a lot of different ways to get that output. So I love that insight, Lori, thank you.

You'd mentioned this in that last response, and I want to jump into it kind of headfirst for our next couple of questions, which is part three of the series that you both put out, titled "The Human in the Loop," which I love this framing, and I would like to unpack that concept for our listeners for a couple of minutes here. Where in the leave and accommodation lifecycle, from your perspectives, is human involvement not just a nice-to-have, but an absolutely essential component, whether we're talking about discretion, oversight, empathy, or legal judgment?

Trish Zuniga: Actually, there are two common oversight models that you'll hear about in the AI world. The first is human in the loop, and that's where humans stay directly in the decision flow. What that means is if the AI generates an output, then a person still reviews it, validates it, and they can still override it before anything becomes final. Humans are still making or confirming the decision.

The second oversight model is human on the loop. And this is where the AI operates more independently day-to-day. Humans aren't reviewing every decision. They're monitoring the system at a higher level, so just stepping in when something looks off, like there's an anomaly or an exception or a performance issue.

Which approach really makes sense depends on your business process. So if you've got considerations like risk level or regulatory requirements, especially in the leave and accommodations field, also ask, what is the impact of getting a decision wrong? These are the

sorts of questions that should guide you whether you need hands-on human involvement or just higher-level human oversight.

And more so in the ADA space, in the accommodation space, human oversight is essential because in the ADA world, we talk about individualized assessment. This is completely contrary to an AI, which provides a computerized generalization. So AI can streamline all of these processes, but the point that we make over and over again is that it shouldn't override human judgment, especially when accommodation requests or disability-related accommodations are involved.

Lori Welty: And one other thing I want to add to that, Trish, that's great. It made me think of a few other areas that I feel like we should call out for human involvement that can be really helpful in these leave and accommodation situations. One is areas where there is employer discretion. So, you know, there's a lot of leave laws out there that allow an employer to waive certain requirements under certain circumstances. And also consider whether or not there should be exceptions to the rule. You know, one really obvious example is where you have requirements for an employee to give notice for a foreseeable leave, but is that leave really foreseeable? And maybe the employee says, "well, it wasn't foreseeable and here's why." That's going to be something where you really want to implement some human discretion there to understand what happened here and should we be waiving this notice requirement in this case. Or, you know, maybe the employee didn't get paperwork back on time. And why was that? And do they have a legitimate reason? They used due diligence, but for some reason they couldn't get it back in time. So there's certain situations like that where an employer has some inherent discretionary rights and obligations under a law that they should exercise.

And then along with that is the idea of empathy. I think completely taking humans out of these incredibly fraught situations is a dangerous idea. You know, employees are coming to their HR with very emotional reasons. They've got someone in their family who is struggling with a significant health issue, or maybe they are struggling with a significant health issue, or maybe they've had a pregnancy with an adverse consequence. These are things where we really should not be fully outsourcing these to machines. We need to have human empathy to build trust and reduce misunderstandings and really to prevent disputes, if I'm looking at this from a compliance standpoint.

And then finally, just in terms of keeping that human in the loop, is just the idea of accountability. And we talked about this earlier, but just that there is going to be somebody held accountable for incorrect decisions. And so it's so important to have those humans in the loop so that we know who made this decision, how are they justifying it? And maybe it was a good decision, but we need to understand that they had that accountability so that they were prepared to explain what happened and why.

Josh Seidman: Thank you both for that. Really, really good overview outline. I wanna piggyback on that last response. Can you share a specific example or two, maybe something that you've seen in the field or that's come up in your work, where human oversight caught something that AI either got wrong or could have gone wrong if it was not overseen and left to its own devices? I think our listeners would really benefit from a cautionary tale or two if you have any.

Lori Welty: Well, I'll kick that off, and I'll just refer back to that presentation that we did. It was a very fertile ground for this because we were really trying to provide AI with the opportunity to do that. So we were asking lots of different questions from different perspectives. And one of the

things that I thought was most interesting—we asked it a question about, it was a pretty straightforward FMLA question about how much time an employee had left and what happened when the employee ran out of time. And it introduced a completely believable-sounding legal construct for how an employer should look at employee requests that exceeded the amount of time available. And it had me researching further to say, where does this come from? I never heard of it. And the way that AI phrased it, I think that if you're not a legal professional and maybe you don't even know where to research, you wouldn't know where to go to research that. That would be a very convincing piece of information that the AI provided. So I think that we should recognize that AI can come up with some really convincing-sounding legal frameworks, and just to be cautious what side you're on.

Then the other thing I'll say is—well, I think those, you know, I think they call those AI hallucinations. I love that phrase. I think those are the most dramatic and the most interesting type of errors that AI can make. I don't think they're actually the most common. The most common or dangerous ones—it's where the answer is really almost right, but it's just missing some context. So if we talk about that example of an employee who, let's say, has 12 weeks of leave and they've exhausted their 12 weeks of leave, and the question has been asked of, does this employee have any FMLA leave left? And the answer might come back and say, "nope, they do not have any FMLA leave back." And that's a correct answer, but we're missing a whole big piece of the puzzle. Are they entitled to an accommodation? Was this a leave for their own serious health condition? Does that serious health condition constitute a disability? Do we need to look at ADA? Is there PWFA?

Asking all of those follow-up questions, I think, is something that sometimes AI will do that, but it's not universal that it's gonna think outside the box. So I think the issue isn't so much those dramatic hallucinations, but just making sure that the AI is not just answering the question that you ask. And if it is, basically that you're recognizing that that's what it's doing. So you're using your own human brain and oversight and discretion to explore further.

Trish Zuniga: And I've got a few cautionary tales from recent cases that I wrote about on our blog, but this one is my favorite because it's a story about a CEO that went into their AI platform, fed it information from an ongoing case, fed it information from their acquisition agreement—they wanted strategies on how to bypass the advice of their in-house legal counsel. And so the cautionary tale there is the court obviously ruled that this was against the law. And when we do these presentations at DMEC in front of brokers, we often say what we're saying to you isn't legal advice—please seek the advice of your counsel. Clearly, the intermediate step isn't to feed your situation into an AI platform and then bypass the advice of that legal counsel, because that human legal counsel would have caught on pretty quickly that what you're trying to do is the wrong thing and tell you, don't do it.

Josh Seidman: Thank you both for that last response. Really interesting examples. I have one final question as we wrap up today's episode. I would love a look forward from each of you. If you had to give our listeners—whether they're employers, HR professionals, carriers, TPAs, claims examiners, and so forth—one or two pieces of practical advice for how to responsibly engage with AI in the leave and accommodation space over the next 12 to 24 months, what would it be?

Lori Welty: I can kick that off, Josh. So I think the first thing that I would say is, returning to one of the things we talked about earlier, is really making sure that you're thinking about how you can use AI to help reduce administrative work and not replace professional judgment. So I think helping AI to expedite you to get to the same place that you might have gotten to already, but

get there faster. And I think there's a tremendous value in letting AI handle those tasks that don't require that high level of expertise—you know, summarizing voluminous documents, coming up with first drafts, identifying missing documentation, tracking deadlines, helping to expedite those topics that really are what you would consider administrative. And those are ways that AI can help make people more efficient, but are not replacing human judgment.

And I think that AI can still be really helpful in areas where you do need human judgment. So things like identifying areas of missing documentation. I mean, when an employee submits a certification, there does need to be a decision about whether they have provided complete documentation. And at some point in time, if they haven't, that could be a judgment, that could be a decision that could impact that employee. So helping at one point in the process to say, you know, this AI can analyze what is required, what they have submitted, and whether it meets all those requirements. And then it gives a signal to the employee to say, we've identified this, this may be an issue. And then the employee can go in there and double-check everything, make sure they agree with it. So it can help flag those items, but not actually making those decisions. So to summarize that is, using AI to reduce your administrative work, but not replacing professional judgment.

And then the second thing I will say is that I think—and hopefully every company who's listening to this already has this in place—but making sure that your company has developed, is working on continuing to develop, its AI governance. Parts of that are things like, from a very basic level, understanding what your company ethos is around AI. Some of our clients—every company has its own ethos about this. Some are rapidly accelerating AI, others are holding back on it. Some are really prescriptive about how you can use it. Some of them are more open-ended. Some require AI in quotas, others not at all. So understanding what your company ethos is about that is really important.

Providing that training for the AI fluency, like Trish mentioned, is so critical as well, and really enforcing it. And again, this comes down to ethos, but really enforcing those things like critical thinking and asking, making sure that the AI is giving you a full perspective, and making sure that's part of how you view the AI used in your organization. So I think recognizing every company, everybody in every company is using AI at this point in time, so really making sure you're clear on how you're using AI.

Trish Zuniga: My takeaway starts with an acknowledgement that everybody is going to be using AI—from an employee perspective, from an employer perspective. So the takeaway is really, make sure you have a good handle of how that AI is being used. Is it being done according to a lawful and confidential way? So from an employer perspective, you make the safe option the easiest option. Make sure that the employees aren't using shadow AI and just using consumer-grade AI to help them with their work. From an employee perspective, make sure that you are using enterprise-grade AI that has the proper levels of confidentiality and governance, and make sure that whatever data you're using there, you're not exposing sensitive information or you're creating legal issues for your organization. So use AI responsibly according to the principles of responsible AI and in compliance with what's already out there, being compliant with FMLA and ADA and other leave and accommodation laws.

Lori Welty: Can I add one last thing in, Josh?

Josh Seidman: Absolutely. Of course.

Lori Welty: Thank you. One of the things I think is just a common refrain when we're talking about AI is the fear of how AI is going to replace us, right? And no matter what it is we're doing, I think no matter what our job is, there is this concern about how is AI gonna impact our field. As we're looking at our kids that are growing up, what fields do we tell them to get into so they're not being replaced by AI? But I think through this conversation, I think it's really clear that AI is not going to replace leave professionals and HR teams. But I do think that there needs to be an acknowledgement that it is going to change what makes them so valuable. The routine administrative tasks that are maybe, you know, they are very important, but they don't require as much discretion or judgment—I think those are gonna be things that are increasingly more automated as we talked about. But the skills like judgment and communication, empathy, navigating those complex situations, I think those are gonna become even more valuable as we accelerate some of the other topics. So I think that there is just as important of a role for leave professionals and disability examiners and HR teams. I just think that the qualities that are going to be more important are going to change in those areas.

Josh Seidman: Trish and Lori, that was so much fun. Thank you both so much for joining us today and sharing your expertise, your insights, the work that you and your FINEOS team are doing to move this conversation forward. Thank you, thank you, thank you.

Trish Zuniga: Thank you for having us.

Lori Welty: Yeah, we really enjoyed it. Thanks so much.

Josh Seidman: Me as well. And thank you to our listeners for tuning in for today's 50th episode of Take It or Leave It. We will see you next time.