

The Employment Law Counselor

Episode 21

PLUS STAFF: [00:00:00] Welcome to this PLUS Podcast, *The Employment Law Counselor*. Before we get started, we'd like to remind everyone that the information and opinions expressed by our speakers today are their own and do not necessarily represent the views of their employers or of PLUS. The contents of these materials may not be relied upon as legal advice.

Victoria Fuller: This podcast is a collaboration between White and Williams LLP and the Professional Liability Underwriting Society, commonly known as PLUS. If you like what you hear, please give us a five-star review, tell your friends [and] subscribe so you never have to miss an episode.

Hi, everyone. Welcome back to *The Employment Law Counselor Podcast*, we're your hosts, Victoria Fuller and Laura Corvo. Today we're doing an update episode on the use of AI in the workplace. Now, we covered this topic [00:01:00] last year, but it's actually pretty amazing to think about how much has changed in just a year.

Laura Corvo: That's so true, Vicki. Quite a bit has changed since we last visited this topic, and there's a lot to discuss. In today's episode, we are going to try to introduce you to some of the different kinds of artificial intelligence that we colloquially refer to as AI, discuss the different ways that our employers are incorporating this AI into their workplace, and then discuss how AI is presenting employers with exciting new opportunities, but also exposing them to novel legal risks. Before we get started, let me introduce today's guest, Victoria Ranieri. Hey, Victoria.

Victoria Ranieri: Hey, Laura. Hey, Vicki.

Laura Corvo: Welcome Victoria. Everyone, Victoria is an associate in White and Williams Boston office. She specializes in employment law as well as higher education law. [00:02:00] Victoria, Vicki and I all work together quite a bit, and I suspect that we're all going to have a few stories to share today. Again, Victoria, we are really glad to have you with us

Victoria Ranieri: I am so excited to be here. As you both know, I've been nerding out about AI so much lately. I'm excited to use our collective brain trust of human intelligence to dissect some artificial intelligence.

Victoria Fuller: Yeah, I'm excited to talk about this too. And actually, Victoria and I have a case where this actually already has come up for us, so we have some perspective on it.

But before we jump into the details, let's start with an overview of what is AI. As Laura said, AI is not one thing. It's actually several different types of technologies, and they operate in different ways. They have different applications, but essentially, they're programs that simulate human intelligence processes, such as learning reasoning, problem solving, understanding natural language, and perceiving and interacting with the environment. Now, a key part of that definition is it [00:03:00] simulates human intelligence; it does not substitute for human intelligence. That's a very important distinction that we'll revisit again and again today.

One thing that I actually learned about AI in preparation for this podcast was some of this technology, the creators don't even understand how it works. It's so fascinating. So again as we're talking about this and we're talking about the problems caused by AI, just keep in mind that, we're not talking about a brain, we're talking about a machine. Let's talk about the different modes or different types of AI. Some of these you've probably have heard about already; you probably have heard of actually several of these, but let's walk through them.

The first is machine learning. That's the development of algorithms and statistical models that enable computers to improve their performance on a specific task by learning from data. So generative AI, which is one you've almost certainly heard of, that's a [00:04:00] form of machine learning, which is particularly relevant to our discussion today. Generative AI is a type of AI that creates a novel output in response to a prompt. That output is generated based on the data that the algorithm was trained on. We'll talk about that more in a minute.

The next type is deep learning. That's the use of artificial neural networks to analyze and model complex patterns and very large data sets. This is your Siri, your Alexa, your autopilot, your Netflix recommendations. It's also used to analyze large amounts of data to detect fraud. So again, many different types of applications there.

Natural language processing enables machines to understand, interpret, and generate human language, as well as to respond. This technology makes applications like chat bots, language translation, and things like that possible; autocorrect, for example. I like to call it auto incorrect because it's so frequently wrong. And predictive [00:05:00] text.

The next one is expert systems. Those are AI systems designed to mimic the decision making abilities of a human expert in a particular area using knowledge and reasoning to solve problems. For example, this is like a medical diagnostic program in which you input symptoms and the program will follow the pre-programmed rules to render a diagnosis or a result. This can also be used in other applications. For example, a tutoring system that tailors the program to that particular student's learning capabilities so each time the student answers a question, it modifies what comes next.

The last one we're going to talk about today, although this is not the only other type of AI [it's] just one that we're going to mention, is agentic AI. That's a system designed to act autonomously to achieve specific goals, making decisions, and taking actions without constant human guidance. This is actually really cool. So, this technology could, for example, monitor traffic and weather patterns and determine an alternate [00:06:00] route in the human resources world. It can generate job descriptions, job postings, screen resumes, and do other tasks like that.

Laura Corvo: So, Vicki, now that we've kind of reviewed the different modes of AI and we see all the different possibilities that AI can present, Victoria, I was wondering if you could tell us how we're seeing employers incorporate AI technologies and tools into their workplace.

Victoria Ranieri: Yeah, totally. For the most part, employers seem really excited about AI. They see the use of these new technologies as a way to increase productivity, increase sales, cut costs, improve customer relations, and create faster response times. Employers are now using AI to do almost everything that human workers can do, which makes a lot of people nervous, from content generation to data analysis, to customer service, to diagnostics, and even agenda setting and note taking.

Victoria Fuller: We're also starting to see employers use various AI tools to [00:07:00] manage their employees and take on what we're traditionally human resources tasks like recruiting and hiring, timekeeping, compensation and evaluations. Victoria, will you walk us through some of those different types of technologies that are available to employers and the human resources management task subset?

Victoria Ranieri: Absolutely, let's start with hiring. Employers can use applicant tracking systems. They use machine learning and natural language processing, and they automate the recruitment process. For example, they may

use keywords to screen resumes and highlight ones that have the correct keywords that match a job description.

After the application tracking system has done its magic, once an employer has that pool of preferred candidates, a next step typically might be an interview. And here again, AI can enter, particularly if the interview takes place over a video powered platform. Employers can also use AI to take notes of the interview, and later they can [00:08:00] summarize those notes maybe for other decision makers who weren't in the room.

Now, if a candidate does well, the next step might be a background check. AI is getting involved here as well. Using machine learning, employers can use AI to perform a background check by gathering data from a variety of sources, including not only public criminal records and employment history databases but also education records and public social media profiles and other system checks.

Let's say you have a candidate who you've selected, you want to make an offer. The next step might be to determine compensation. AI can help you there too. They can do your market research for you. Once the employee comes on board, both the employer and the employees can use AI in a lot of ways in the course of their jobs.

[For] HR, specifically, AI can be helpful with performance evaluation systems. They can track metrics. They can also evaluate recorded customer interactions. [For example], every time you call customer service and they tell you the call's being recorded. I never believed that [00:09:00] anybody was actually listening to my calls, but now that AI is available, I'm starting to second guess that belief.

AI can also be used in timekeeping systems. Lots of employers are switching to biometric data now in their timekeeping systems.

Victoria Fuller: I think these tools create some pretty exciting opportunities. They have the potential to create efficiency, cost cutting, but it seems to me like they can kind of be a double-edged sword because they also present legal risks, right, for employers? We're starting to see a pretty quickly evolving regulatory and statutory landscape. Laura, can you walk our listeners through what we're seeing on that front?

Laura Corvo: Yeah, sure, Vicki. The double-edged sword for employers is that the algorithms that run these new tools that Victoria just described, may in some instances have a propensity to obtain results which generate a disparate impact

for certain protected classes, right?[00:10:00] Intuitively we think, well, we're taking human biases out of the process and we're leading it to the machines. But the underlying algorithm that runs the machines, it has been found in some cases, has the propensity to discriminate against individuals based upon a particular race, gender, age, or disability. We're seeing regulation start to crop up, which are designed to prevent this type of discrimination.

Let me first just talk about the federal landscape and where things are at because we've seen certainly a big switch from the last time we did this presentation. Currently, there are no federal laws regulating AI use in the employment context. Under the previous administration, there was a pretty comprehensive executive order, which called on agencies like the EEOC and the Department of Labor to regulate AI and specifically to address the potential disparate impact that may result from certain algorithms. Under the previous administration, the [00:11:00] EEOC, for example, issued guidance that said, "Hey, if these AI tools are disproportionately affecting protected groups, that's a violation of federal discrimination laws like Title VII." Well, that guidance is now withdrawn.

There's recently been executive orders from the current administration, which are kind of aimed at reducing federal oversight of AI, taking a step back from regulations in this area. There was also an attempt in Congress during the passage of the Big Beautiful Bill to attach a provision which would say that states could not regulate AI for at least 10 years. That provision did not pass.

It's really at the state and local level where we've seen a lot of action and a lot of regulation in this area. I'm not going to go into tremendous detail, but let me just highlight a couple of the state and local laws that are in play because I think we're starting to see a pattern and I expect we'll see more.

New York [00:12:00] City was the first to regulate this area. They passed an ordinance in 2021 that took effect in 2023 that prohibits employers from using automated employment decision making tools, many of the tools that Victoria described earlier, to make employment decisions unless the employer conducts a specific bias audit to make sure that there's no propensity for that discriminatory result. And it requires the employer to post a notice to employees and job applicants that they're using the tool. [This is] to kind of put them on notice of this so that they know this tool is being implemented. If they perceive that they may have a discriminatory impact as a result of it, they can take action.

The New York City law, even though it's just New York City, it has a pretty broad reach. It applies to any job in New York City, even if the job is only part-

time in New York City. It applies to any employer who is based in New York City who's using these tools. It also applies to any [00:13:00] remote job that reports into a location in New York City.

California recently passed a regulation with their Fair Employment and Housing Act. This regulation addresses how there could be discriminatory impact through the use of this tool. It requires employers to keep records for at least four years. While the tool doesn't require employers to specifically conduct a bias audit, it does recommend it to prevent and as a defense in discrimination claims.

Colorado is another state that has passed a very comprehensive AI law. And again, I'm just not going to get into everything. For employers, it basically means you have to use reasonable care to protect against what is foreseeable algorithmic discrimination. That includes, and the [00:14:00] regulation gets pretty specific, having a risk management plan, completing an annual impact assessment or a bias audit, providing notice to employees, explaining to employees who have an adverse decision that this tool was used and giving them the reason why they were rejected from the job, and an opportunity to appeal that decision. [It] also [includes] putting a disclosure on your website that you're using these tools.

The Colorado law was set to go into effect early next year in February, but I think there's been some pushback about the impact of the law and whether or not some portions have to be scaled back. They've kind of kicked the can down the road until June of next year. So kind of put a pin in this one and we'll see where that goes.

Victoria Ranieri: Laura, I was listening to some really interesting reporting on that. Vicki was saying at the beginning of the podcast how some of the technology is outpacing even what people know about it. [00:15:00] I was listening to some reporting on the special session where they determined in Colorado that the law would be pushed back four months. The reporters were saying that some of the difficulty in the conversations between the lobbyists and the legislators was the understanding of the terminology in the law because it is just so far outpacing what we really understand about it as, you know, people who aren't in the industry. So I just thought that was a really interesting tie in there.

Laura Corvo: Yeah, it's a really interesting point. I think that it's going to be difficult for the regulators to get ahead of the technology. The technology's moving really fast for employers. You're both going to have to keep up with the

technology and also understand regulations, which may not make sense in the environment in which they're being implemented. So, it's a bit of a wild west type scenario for employers in this area.

There's other states who pass regulations Illinois [is] another one that amended its discrimination law to say, [00:16:00] “Hey, if you commit this algorithmic discrimination, it violates our discrimination laws.”

Even in those states where there haven't been regulations or amendments passed to discrimination laws, the state agencies that regulate the laws have issued guidance. For example, Massachusetts Attorney General says that anti-discrimination laws apply to AI. So, you've really got to be careful in this area.

Victoria Ranieri: Wild West was a good way to word it. It's just changing so fast right now. Before the podcast, I actually asked AI itself to sum up the status of AI regulation in 2025. I didn't want a comprehensive summary, so I said, “Do it with some humor, please.” The tool I used to sum it up said, “Please ensure your model doesn't hallucinate, discriminate, manipulate, or accidentally run for office.

As far as I know, AI hasn't run for office, but several lawsuits have arisen in the employment context, accusing some AI applications of discriminating against employees or potential [00:17:00] employees.

And this is a real emerging risk for employers, right, Vicky?

Victoria Fuller: Yeah, it is. I'm going to compare and contrast a couple of recent employment cases involving AI. Before I do that, I just want to put it out there. These are just two cases to give examples. There are many cases already in suit driven by the use or misuse of AI.

An issue that we see coming up repeatedly are individuals using AI not as a tool, but as a replacement for human intelligence or for human work product. Again, I know I said this before, AI mimics human intelligence. It is not a substitute for human intelligence.

The other thing that I want to just point out is something that we saw repeatedly in case law, and this dovetails back to what Laura was talking about with these regulations, is cases where it appears that the defendants didn't understand or were not aware of the law. I'll just give some examples, this is not in the employment context. [00:18:00] [There are] cases where companies using AI, for example, were violating the BIPA statute in Illinois because they were using

AI to capture biometric data about customers, for example, and not giving the required notifications and getting the required consent from customers.

My point there from an employer's perspective is if you are using AI already in your business and you have not gotten with counsel, please do. Particularly if you're a multi-state employer, you got to get with counsel and make sure you're complying with the law in every state in which you operate. As both Laura and Victoria said, the regulatory landscape is changing pretty rapidly and you don't want to be behind the ball on that.

Okay, so now let's talk about these two cases. The first one is iTutor Group. That's *EEOC v. iTutor Group, Inc.* This was a case filed in the Eastern District of New York. The defendants in that case operated an online tutoring [00:19:00] organization that provided English language tutoring services to students in China. The organization used a program where they actually had programmed it to reject applicants over a certain age. It was like 55 for women and 60 for men. The program would look at the resumes and--- oh, what I should mention, that the online application actually solicited ages for the applicants, which is a huge red flag--, but anyway, the program would see if you're over 55 and female, eh, you're out. If you're over 60 and you're male, eh, you're out. The program actually ended up resulting in the rejection of approximately 200 applicants who were over that restricted age. iTutor ultimately ended up settling with the EEOC for approximately \$365,000.

Now, this is a pretty straightforward case, right? Obviously, we're looking at this. This looks like intentional age discrimination because they program the software to kick [00:20:00] out applicants who are over a certain age. Sophisticated AI programs, though, can inadvertently use proxies for protected categories. For example, they may look at the number of years of experience and without the employer intending to do so, the program might look at that and start calling out resumes basically with too much experience, which would be a proxy for age discrimination. We know these kinds of things are happening, whether it's age discrimination, race discrimination, other types of discrimination.

The AI programs can take that data set. They look at what's been used to feed them and where it doesn't match up. [They] call out resumes based on not necessarily a protected category, but a proxy for a protected category like individuals who went to historically black colleges or women's colleges because those were not part of the sample set.

Victoria Ranieri: That's way more insidious too, right-- like the iTutor case that you were talking about. If a human did that, had somebody [00:21:00] walk into an interview said, "How old are you?" and when they said over 40 told them, "Goodbye," we would know how bad that is. That's what the algorithm was doing. But these examples that you're giving are much more difficult to know that the algorithm's doing that and to correct it, right?

Victoria Fuller: Yeah. But also keep in mind as well that these are applicants, right? It's not one single employee. These are applicants. So, you could have potentially a very large class of plaintiffs. In iTutor it was 200 people. You have a large employer, say you have 500 applicants for a position [and] half of them are over the protected age category and they get screened out, your potential class action is just ripe for the taking right there.

Let's talk about a different case to give a different example here, which is tying into what I was just talking about with this more nuanced discrimination. *Mobley v. Workday, Inc.*, that's a case from the Northern District of California. This is [00:22:00] one where the plaintiffs allege that the Workday's AI algorithm discriminated against employees and did it on the basis of age and disability and race, so we have three protected categories. That case was filed last year. And again, this is a case where we have a very large potential punitive class because Workday, Inc. itself is so huge.

Just to give background on Workday, Inc. provides human resources management services to medium and large size employers. It uses an applicant tracking software, which is what Victoria was talking about earlier. Workday uses that software to screen resumes.

The plaintiffs in that case alleged that the algorithm had a disparate impact on African Americans, older and disabled applicants who were screened out at significant rates from applicants, not in those protected categories. They actually had some statistical data in the complaint, which I thought was very interesting. The [00:23:00] complaint also says that "the algorithms too often have discriminatory intent, even where demographic data such as race, age, and disability are not included as inputs." This is because algorithms can learn to use omitted demographic features by combining other inputs that are correlated with race or another protected classification like zip code, college attendance and membership in certain groups.

This is the part that would keep me up at night as an employer using this type of software because you don't know what the algorithm is doing necessarily, right? You don't know if it's picking up on things and then creating a disparate impact

on certain protective categories of applicants because the algorithm says, “Oh, well I don't see anybody in my feeding pool who went to a women's college, so we're going to exclude out all the applicants who went to women's colleges.” That's what this case is really about. This is not the clear-cut iTutor case; this is the case where the [00:24:00] algorithm is alleged to have discriminated based on these proxies, so to speak.

Again, as I mentioned, this is a putative class action and Workday, Inc. is a very large company. According to the complaint, as of April, 2023, nearly one in four of all US openings were processed on the Workday platform. Just let that sit for a minute.

Victoria Ranieri: That's really scary.

Victoria Fuller: Yeah, it's a big potential class. There's something else that we're seeing with these AI cases, these class actions. Keep in mind, again, because these are applicants, they don't have to show that they were ultimately qualified, that they would've gotten the job. They just have to show that they were denied the opportunity to compete for that job. That's a pretty big, again, potential class in this case.

There's another interesting allegation in this complaint that I want to discuss. The complaint alleges that part of the application process required by the named plaintiff was to take a personality test, [00:25:00] which was not consistent with any business necessity. The complaint alleges that that personality test was designed to identify mental health disorders or cognitive impairments and therefore create a disparate impact on disabled applicants who are screened out by the test.

Now, if you're a listener of this podcast, or just otherwise immersed in the employment law world, we all know that mental health issues are a hot topic for the employment world. It is a hot topic for discrimination law. I think this is something else here that employers want to be thinking about. Should we be using these personality tests, whether it's AI related or not? Is that creating a disparate impact on individuals with certain either mental health conditions or neuro divergencies?

Laura Corvo: Vicky, hearing the two cases you just described, especially the Workday, Inc. case, I'm sure a lot of employers are concerned about class actions or potentially very expensive litigation. I'm sure some of our listeners are wondering what [00:26:00] employers who do want to use some of these AI tools, like Victoria described earlier, for their human resource functions to do

things like help sort out resumes or the like, what should those employers be doing to avoid being in a situation that Workday, Inc. was in?

Victoria Fuller: That's a really good question. The answer is you do have to periodically check the output to ensure the algorithm is inadvertently creating a disparate impact or discrimination against a protective category. Again, this is one of those features of some of these laws that we see coming into effect as well. What to me I think is interesting though is, functionally, if you're having to check the output, it may end up making the use of that algorithm in the first place not even worth it if it doesn't save you any time because you're spending time trying to make sure that it's not discriminating.

It may also, I think, be difficult to verify the results. Say we have an algorithm that's using experience as a proxy for age. How are we going to [00:27:00] verify what the age of all the applicants are? Are we going to have to go through the resumes and a sample set to manually add it all up to see what the number of years of experience are and try to estimate whether those candidates are in the protected age category? Are we going to come up with a sample set to feed into the algorithm and see what it does with it?

I'm not an expert on how you test these things. I just would be concerned [that] even if you're trying to verify the results, it may not be able to truly verify what's going on in the ghost in the machine.

Laura Corvo: Vicki, since many of the state regulations we talked about earlier have specific requirements on conducting these bias audits, it might make sense to engage counsel to at least ensure that whatever audit or check you're doing complies with these requirements.

To your point, Vicki, employers will have to kind of weigh the costs and times associated with conducting these checks and audits against the value of the tools themselves to see whether it makes sense [00:28:00] for them to proceed with it.

Victoria, in addition to the risks in using the AI to perform human resource functions, we're also seeing risks when employers use AI tools to perform tasks that were traditionally reserved for counsel.

Victoria Ranieri: We definitely are. AI as your lawyer, is a real trend right now and one that can be dangerous. There's a real temptation to use AI as a substitute for counsel, either to ask advice on a situation an employer might be facing or to draft important legal documents. For example, are you onboarding an employee? Well, it's pretty tempting to use AI to draft a non-disclosure

agreement or a confidentiality agreement. Do you have a question about a reasonable accommodation for a potentially disabled employee? [It's] also, tempting to ask AI for help there. What about if you're terminating an employee? It's tempting to use AI to draft a separation or separate agreement.

Look, we get the temptation. AI is super accessible now. At first blush, it may seem more cost efficient [00:29:00] than hiring counsel. The work product looks just good enough that to an untrained eye. It may appear that AI has you covered, but unfortunately, using AI as your lawyer could really cost you in the long run.

To demonstrate this, I did a little experiment. I had AI to draft a separation agreement. Employers use these a lot to have some comfort that a departing employee won't later sue them. Laura, I've heard you call this a piece of insurance for employers before. A really important part of getting that insurance is getting an adequate release so that the employee can't later come back and sue you. To try and do that with AI, I used the following prompt: "I'm a Massachusetts employer. I have to lay off an employee who has worked for me for five years. Give me a template separation agreement. I want to offer them \$5,000 in severance." AI took less than 20 seconds to generate this agreement. It gave me a [00:30:00] reasonable professional looking template, but in reviewing it, it omitted many important things that a competent lawyer would've included.

The first thing is that release that we talked about. I know that we all know to include the company in a release, but generally we also either define the company to include employees and managers, or we include them separately as release parties. AI didn't do that. It didn't ask me if the company had a parent or an affiliate that should be included. These aren't minor things because this means if an AI generated release like that were given to an employee, it wouldn't prevent the terminated employee from suing his or her supervisor personally or suing the parent company. You save the cost and paying a lawyer to draft the agreement, but you may be paying a significant cost of defense.

Second, although AI did flag that the older Workers Benefit Protection Act may apply-- which I have to say I was impressed with, I didn't think it would get [00:31:00] that issue, and I purposely didn't give it the employee's age-- it did say that if the employee were receiving the separation agreement and they were over 40 some OWBPA language needed to be included. It left out a lot of language too. For example, it didn't say that the employee had the right to consult with counsel. It did give the 21-day period to consider the agreement.

But even though I told AI that this was a layoff, it didn't ask me if more than one employee over 40 was being laid off.

Why is that important? If there is a group of two or more employees that are laid off that are over 40, they each get 45 days to consider the agreement. And again, these omissions aren't minor because if this language is applicable and it's not in there, the employee may not have fully released his or her age discrimination claim. So, you've paid them the severance and then they can sue you.

Then there are other issues too, like not containing a no rehire [00:32:00] provision. Vicki, I know that's one that you have a lot of experience with where, sometimes if that's not included, the employee then reapplies and doesn't get hired and then they can sue you.

The point here is this, the AI separation agreement had several issues that potentially left claims unreleased. If I handed this in to either of you as your associate to use for one of your clients, I think you would've graded me like a C minus at best? We always aim for A's here, so I'm not going to be delegating any of my assignments to AI, don't worry.

What you've saved, quote unquote, in not engaging counsel you may well be paying 10 times that in fees to combat the deficiencies in this agreement. So sometimes the same, you get what you pay for really rings true.

Victoria Fuller: Yeah. And it's also true that prophylactic legal services are almost always more cost efficient than defense [00:33:00] costs when you have to respond to a claim. I think lawyers and doctors too have been saying this for ages. The internet is not your lawyer, the internet is not your doctor. AI is not your lawyer. AI is going to miss these things because again, it's not a brain, it's not a human. It's mimicking what a human can do, but it's not a perfect technology.

Actually, I'll give just a quick example. A friend of mine sent me some headshots that she had had AI do recently. It was actually really interesting because it was good, like you looked at it and it was good. But [in] several of the pictures the AI changed her face in ways that, as a human, you would look at it and you knew that that was not what she looked like.

I mean, if you just saw those photos, you wouldn't necessarily know that. It was like, "That's weird. Her face looks different." But that's the thing with AI, it's

not a human. It's not going to know that it changed her face in material ways that made her not look like herself.

Victoria Ranieri: The devil is totally in the details.

Victoria Fuller: That's right. The devil is in the details and the details on some of [00:34:00] those headshots. It was interesting. Like they were good pictures. It just was weird in that they, some of them didn't look like her. So anyway, the point is, it's not your lawyer. Please call your lawyer and have your lawyer do traditional lawyer things.

Okay, so Laura, let's talk about risks. We love to talk about risks. What are the risks that employers are facing when they use ai?

Laura Corvo: One of the biggest risks is confidentiality concerns, right? Again, employers are really wowed by the various things that these AI tools can do, but a lot of times they have not fully considered what happens to the information and data that they put into the AI tool.

As employment lawyers, we know that we take great pains to protect employer confidentiality. We make employees sign confidentiality agreements. We make people sign non-compete agreements. We try to enforce those agreements when our employees break that confidence. We're often trusted with our customers confidential information, and we put up guardrails to [00:35:00] protect that.

But what happens when information is put into an AI tool? Do we know whether or not that confidential information is protected? Is it disseminated? How is it going to be stored?

There was an example, I guess, about a year or so ago where Samsung got caught and wound up banning the use of Chat GPT because one of their engineers uploaded sensitive source code into Chat GPT to check something and, as a result, publicly disseminated that source code. Source code that was proprietary and that was key to their business suddenly gets out there because somebody puts it into an AI tool.

Another area where confidentiality concerns come into play, and I've been working and seen as a lot of clients concerned about this, are with AI note takers and assistants. Many of the virtual platforms we know-- Teams, Zoom, Google Meet, etc-- have AI tools, which can either [00:36:00] record or transcribe or take notes and send follow ups for meetings. On the surface, it sounds like a really great idea, right? We've all been in meetings where nobody

remembers what happened or nobody does follow up and the time spent in that meeting becomes completely unproductive. If you have a tool that can summarize what happens, send a follow up agenda item, it's really attractive. It could be a very valuable tool. But before you use those tools, you have to take a step back and think about how that AI tool is processing and disseminating and storing that information.

Victoria Ranieri: Really quick on that, Laura. I had a creepy experience with AI while preparing for the podcast about it remembering information.

Vicki and I have a number of things coming up that require us to send in short bios, so I figured I would delegate that task to AI on my behalf. I took my longer firm bio and I put it in an AI tool, and I said, "Can you shorten this for me and make it punchy?" [00:37:00] When I was later preparing for the podcast and I said, "Can you tell me a joke involving AI?" the AI tool said, "Sure. Do you want the joke to be about employment law or higher education law?" Which it had remembered from processing and learning my bio.

I do not want AI to know anything about me. I didn't even think that those details would be stored when I delegated that task, so it really is something that [for] employers may not be top of mind until something like that happens.

Laura Corvo: Yeah, and again, we keep going back to this, but AI has a lot of capabilities, but it's not human. It doesn't have the ability to exercise discretion or know what's relevant, but it does have a lot of this information stored. How it spits it out might not be the way we want it to be spit out.

The Wall Street Journal ran an article on the AI note taking functions and [00:38:00] noted that a lot of times what happens is at the beginning of the meeting you have the conversation: How was your weekend? What'd you have for lunch? Suddenly there's bullet points on that as opposed to what's important in the meeting.

The AI tool is, it's a tool, right? It doesn't have discretion that maybe somebody else taking notes would know not to sit there and say, "Vicki ate tuna salad for lunch and Victoria had the ham and cheese." [That's] something that's completely extraneous.

Victoria Ranieri: I don't know. That might be the most important part of some meeting.

Laura Corvo: The dissemination of these tools, how these note takers disseminate information, is also something you have to get your hands on as an employer.

There was an example, Alex Bilzerian, who is an engineer and investor in AI, went on social media, there were a bunch of articles about this. He had a virtual call with employees from a venture capital firm. They have a discussion and he leaves the meeting. [00:39:00] The employees of the VC firm continue to talk, kind of thinking they're talking amongst themselves, and they're discussing very confidential information, sensitive information. And guess what? The AI tool at the end of that meeting disseminates a recording of that meeting, not just to the employees in the VC firm, but to Alex Zoran who now has all this company's confidential information and sensitive information. Now, fortunately, he doesn't do anything with that, but he does decide not to invest in that company because of how they treated their information so lightly.

I think the takeaway from that is that employers who are using these note taking functions have to make sure that they have appropriate settings to make sure that they're not disseminated automatically to any participant in the meeting. Maybe they're just disseminated to the meeting leader, who then uses that discretion in human oversight to just send it to the people who actually need to see it.

I [00:40:00] think there's also concern about sensitive topics, right? We probably don't want the AI tools discussed when we're having an attorney-client meeting. We have to remember that some of these tools, now that everyone knows they're out there, when we have discover requests, first thing that plaintiff's attorneys are going to ask for is, "send me all the transcripts or the recordings of every meeting you had."

Employees have to know that. The same way we train employees that "Don't write anything in an email that could come back to haunt you and be discoverable or be exhibit A to a lawsuit." You also have to do that if you're allowing your employees to record meetings or transcribe meetings. So, all of this is going to come into play.

Victoria Fuller: Yeah, on that point, remember back when email was new, people were writing things in email and it was just this wealth of evidence. Then people learned, "Okay, don't put that stuff in email." They continue to put it in text messaging and then they learn, "Okay, don't put this in text messaging." [00:41:00] Now they're putting it in the work chats and it's like AI is the next thing.

They're not going to think about it during a meeting that's being transcribed, so that's going to be the next place where, "That's where the evidence is, right?" They're not thinking about "This is creating a record." that's where people are gonna speak extemporaneously,

Laura Corvo: At least with the email and the chat you're physically writing. The AI tool is kind of, even though there might be a notification at the beginning of the meeting that it's turned on, it is invisible. It's kind of in the background and people may forget and let their guard down. You've got to make employees aware that, "Hey, this is happening. When the tool is being used, be aware and don't say anything that could lead or compromise you or the company."

Victoria Fuller: Absolutely, Laura. On this note, with all of these risks that we're talking about, do you think employers should not use the note taking function at all?

Laura Corvo: I don't mean to suggest that note takers should never be used. For many employers there are pretty significant value to them. That efficiency and [00:42:00] follow up that you know might be lost without them is certainly valuable, but employers might need to kind of proceed with caution and really understand how they work. Set the appropriate privacy and dissemination controls just to make sure they're not exposing them to this risk.

It's like when I go to the beach. They have the different flags set up. You know, it's not a green flag; you can go in the water and do whatever you want. There's pretty significant waves out there. I'm not saying never go into the water, but proceed with caution. It's the yellow flag. Make sure you are operating knowing how this is occurring and that things could come back to compromise you.

On that, another risk or yellow flag is your brand integrity. When you're inviting employees to use these AI tools, not just in the employment human resource context, but in any context, you have to remember, we said it so many times today, they're not human and they're not perfect.

AI tools are prone to [00:43:00] hallucinations. They can give false, misleading, or non-existent information. We see this in the legal field. You know, somebody writes a brief and cites a case. They run the brief through AI, a case is cited [and] the case doesn't in fact exist. The judge asks for the copy of the case, and it's nowhere to be found.

There's also deepfakes. That's where someone's image or voice is used to manipulate and place[d] in a scenario that doesn't actually exist. We see this in celebrity and politician [cases]. There was [the] example [of] the president of Ukraine, President Zelensky. There was a deepfake of him telling Ukrainian soldiers to lay down their arms and surrender to Russia. Of course, it wasn't true, but it looked very realistic.

AI can also be prone to just some sloppy mistakes, like spelling and grammar errors, things that you don't really want to compromise your brand integrity. If you're going to be using these products, you've got to put a human check over it to make sure [00:44:00] that there's verification for accuracy and other brand integrity because you don't want that egg on your face.

Victoria Fuller: Yeah, we could actually talk about the different risks from AI for the next three hours, but I want to talk about a completely different direction of risk for a minute. And that's the risk from employees themselves to employers.

Employees are using AI to educate themselves on different laws [and] different potential causes of action against an employer. They're using it to help them craft demands to employers, to craft emails to employers either making claims or setting up for, whether it's, "I need an accommodation, or I'm a whistleblower." They're using these tools to help them craft a better email, to craft a better demand letter to send to the employer.

They're also using it in other ways too. You have pro se litigants who are using it to prepare their brief for them. As you [00:45:00] said, Laura, one of the things we are seeing a lot in the legal field are these hallucinated cases where they're either completely fake, or sometimes they're real cases, but they don't say what the AI cites them for. We're finding that courts are as intolerant of pro se doing this as they are of attorneys. Frankly, it's happened so many times. I'm shocked that it continues to happen, but apparently this is where we are. It's definitely making those claims more of a headache because, again, it's an imperfect system, but it is enabling pro se to create a better legal product, so to speak.

I'm not saying it's a perfect product because, again, it is very imperfect for all the reasons we have talked about, but certainly better than what most pro se could put together on their own.

Victoria Ranieri: But the courts are losing patience with these types of products. Like you said, they're recognizing what a huge problem this is.

They're not only losing patients with lawyers, but with [00:46:00] those self-represented pro se parties.

In a recent case in Missouri called *Kruse v. Karlen*. There was a pro se litigant whose appeal was completely dismissed because he filed a brief with 22 AI citations that were incorrect. He was also fined \$10,000 to compensate the opposing party's lawyer for the time spent responding to the cases.

This should not only give pro se pause, but it should give employers some pause too. If you get a brief that is filled with AI or a complaint that looks to be filled with AI, don't combat AI with AI. Two wrongs don't make a right. Two AIs don't make a good brief. Here, really engage counsel. Let them deal with it because it is really, really common and the pro se may not know better, but you do.

I had a case recently where there was a pro se litigant and they admitted to searching every AI database to [00:47:00] come up with their complaint, including Open AI, Chat GPT, Claude 3.5, Google Gemini, Microsoft Copilot, Google Bard, and Thompson Reuters' CoCounsel. The pro se, came to the conclusion that this was a first case of first impression in the history of the republic since 1776. Think you want to guess what kind of case it was?

Victoria Fuller: Breach of contract.

Victoria Ranieri: Exactly! You got it in one. We've never seen one of those before. It is really difficult to know where you're going wrong with AI, so if you get a pro se with an AI claim, still get counsel.

Laura Corvo: And not just the pro se, as Vicki said earlier, it's the employees who are going to write you those notes saying, "Oh, your sick leave policy is completely wrong, and here's a citation," And then you're suddenly scratching your head and changing something. Before you make that change, let's talk to counsel and make sure that employees got it right or that whatever AI source they use got it right because [most] likely it's wrong.

Victoria Fuller: All right, let's change direction completely. Another way in which AI presents [00:48:00] risks for employers. Now we're going to talk about cybersecurity.

I cannot emphasize enough how important this is. Please pay very close attention to what we're about to talk about. AI is enabling threat actors to craft a better threat. Remember 10-15 years ago, the emails would be full of

grammatical errors? You could tell that whoever wrote it didn't speak like native English, or they had errors in them that made them pretty obvious to spot? Now they're running their emails through AI to craft a better sentence, not just in good English or whatever language it's supposed to be in, but also that it's in like business language. It's something that you would expect to say. Maybe it'll say, "We'll circle back around," or use corporate lingo, "We're going to synergize," whatever. You know what I'm trying to say. It'll look professional.

It's not just written texts that AI can [00:49:00] help with, it's also voice spoofing. This is also a big problem. They can mimic somebody's voice. They can call the office and throw the voice through a masked line and instruct a subordinate to transfer money outside of the organization.

Employees must, must, must be trained on this, to be aware of these types of attacks[and] to know what to do if they're not certain about whether something is true or not. It is getting much like you look at a deep fake. [For] some of these you can't tell. You look at it and maybe you know it's not real, but you look at it and it looks like it could be real. That's the kind of thing that we're seeing with these cyber-attacks, so employees should absolutely be trained on it.

The employer should get with their broker and see if cyber insurance is appropriate for their business. One thing that we see a lot is smaller, mid-sized businesses think "I am too small to be a target, so I don't need insurance because no one's going to come after my business." No, [00:50:00] wrong. You're the perfect target because you are small. You don't have the protocols in place [and] you don't know what to do. You are easier to get at. Small and mid-sized businesses are the least likely to get cyber insurance, and they are also the most likely to have these sort of "bet the company" attacks because they're not prepared for them and they don't have insurance to back them up.

Please make sure, again, that you both get insured on that. Also make sure you're getting appropriate training for your employees and that you have policies in place to ensure your data is protected and that employees have a place to go to follow protocol if they don't know what to do.

Let's change directions again entirely. Victoria, you want to talk about IP issues?

Victoria Ranieri: Yeah, there are so many claims related to both input into AI systems and output from those systems. What's really difficult about these claims is the law surrounding them is totally unsettled, and that makes things very unsettling for employers.

[00:51:00] Laura, you talked about things that you may not be aware you're inputting into AI. You know, those meetings where you're catching the conversation about lunch. Well, a lot of the IP cases sort of circle around things that are purposely input into AI to train it. The most famous one was the first one. It's *Thompson Reuters v Ross [Intelligence]*. Thomson Reuters makes Westlaw, which we're all familiar with. It's a legal research tool, and Ross was a new competitor. They made a legal search engine that was powered by AI and decided to train it on Ross Westlaw's copyrighted head notes and the case summaries. You know, those things that appear before the text of each case on Westlaw. Thomson Reuters sued for copyright infringement, and Ross argued that this is a fair use of the material and that it was transformed and that it was not copyright infringement.

When I say the law is unsettled here, I am not exaggerating. The judge in 2023 denied the motion for summary judgment that Thomson Reuters filed. Then in February of this year, he reversed himself [00:52:00] and he allowed much of the motion and said that Ross had committed copyright infringement and the AI training on these portions of Westlaw was not fair use. The case is now on appeal to the third circuit.

There are two cases in California that reached opposite results. They're saying that AI training might be fair use. In *Bartz v. Anthropic PBC*, several authors sued for use of both pirated and non-pirated copyrighted books to train Claude AI. The court held that the use of the books to train AI was, quote unquote, "exceedingly transformative and thus fair use."

Out of the same court came *Kadrey v. Meta*. Authors there sued meta for using fictional works to train LLaMa, its AI platform. In that case, the court took a much more cautious approach, stating that transformation alone does not guarantee fair use, which was not what the judge had said in the previous Bartz case. [00:53:00] But in the Meta case, the judge ruled that since the plaintiffs hadn't proved any harm resulting from the use of their work, like devaluing the work, for example, the court found that Meta was entitled to summary judgment because the training there was fair use. However, the court heavily suggested that had the authors developed a record and established evidence of harm, he would've decided the case the other way.

If you're thinking, "I am an employer. I don't develop AI. I just use it, so I'm not training the tool. I'm safe, right?" Nope. Outputs can also potentially result in IP infringement. In *Dow Jones & Company, Inc. v. Perplexity AI, Inc.*, the plaintiffs there included the Wall Street Journal and the New York Times, and they argued that the output of Perplexity's AI platform infringes on copyrights

because the answers it provides to users' questions include full or partial verbatim reproductions of plaintiff's news analysis and opinion articles.[00:54:00] Other times, the complaint alleges, Perplexity turned copyrighted articles or work into paraphrased versions or provided summaries of those copyrighted works. But those summaries, argued the plaintiffs, serve as substitutes for accessing the works through plaintiff's own websites.

If your employees are using AI tools that are prone to these types of outputs that may be plagiarizing other works, you could be unwittingly using or disseminating copyright material. The Dow Jones case remains pending, so here too, we don't have an answer. This is an area where we are really watching the law because these risks are emerging and we do not understand really how courts are going to treat them going forward.

Laura Corvo: Victoria, it's probably makes sense, knowing that there's this potential risk for copyright infringement, that we have that human check over whatever's going out every time we're getting any output from AI.

Victoria Ranieri: [00:55:00] Absolutely. Even with the really poor AI joke I told at the beginning of the podcast, that AI wrote for me, I did run that through several searches to make sure I wasn't pirating the that piece of humor from some other source.

Victoria Fuller:

Actually, that's a great point, right? You don't know if the output of that AI could potentially be infringing. Victoria, you and I talked about this too. Say you're coming up with a tagline. What if it's giving you a tagline that has been inputted into it because it was part of the data that it was fed with? I just would be very cautious to make sure that anything it kicks out to you is not infringing on somebody else's trademark, copyright, any other form of intellectual property.

Laura Corvo: And remember, AI is not human. [It] can't have novel ideas or thoughts, right? It's only processing what's already out there in the world, so chances are it's stealing somebody else's thoughts or information. You've got to be careful that however it packages it [00:56:00] to you is a way that's not going to present a copyright infringement.

Victoria Fuller: Exactly. We've just talked about several different risks that take on entirely different areas of the law and different approaches. Now let's talk about our favorite topic.

What should the employer do to minimize those risks? Laura, take it away. What do you want to do?

Laura Corvo: All right. Well, first thing we must do is develop a very robust AI use policy, right? That policy is going to require some real careful consideration and not just partnership with your HR folks, but also partnership with your IT team. They're the ones who actually know how these things work.

Each employer's policy is going to vary a little bit. I don't see, at least not in the present climate we're in, there being a one size fits all AI policy for employers that can just be spit out. To that end, please do not use AI to write your policy.

Victoria Fuller: Can I tell a story based on that?

[00:57:00] There was a recent case, I think it was out of Minnesota, where an expert who was an expert on AI used AI to help him prepare his expert opinion. It hallucinated citations, which then got caught in the litigation. It's mind blowing, the expert on AI used AI and it created the exact problem that everybody knows AI creates, it hallucinated citations. If I remember that case correctly, the expert's report got excluded.

Anyway, yes, to your point. Do not use AI for this. Use your lawyer.

Laura Corvo: Yes, and your counsel and your IT people. It's a collaborative effort and there's a lot of thought that goes into these policies.

The first thing you want to do is think about what tools you want your employees to use. Again, this is where you got to get with your IT folks to know how they operate, know what the confidentiality and the dissemination settings are. Know that your list may change over time, so you may want to have something that's evolving. An AI policy that [00:58:00] just gives employees carte blanche to use any AI tool is probably a very bad idea.

Given some of the privacy and other concerns we discussed, you may want to prevent employees from using publicly available tools or tools on their own private accounts. It probably makes sense to invest in a corporate account with a corporate license for the different AI tool where you're controlling the privacy setting, similar to what you do with other software. Make employees only use the software you provide and not things from outside the organization.

Overall, the policy should put guidelines around what is being put into AI tools, things like confidential information [and] checks around what's coming out of

ai. As I said before, we always want those quality control checks, those human checks, whenever there's an output. The policies kind of set the scope of who, when, and for what purpose employees are using AI. You might want to allow a manager to use AI in one setting, but not line level employees or, vice versa. There may be specific types of employees who [00:59:00] are using it, while others are not. You also want to make sure your employees are required to be transparent in their use of ai, that they're not passing off work as their own, which is really ai. That could get us, you know, in copyright infringement or create confidentiality issues.

You want to be clear that violations of the policy are going to result in penalties and likely termination so that employees take these policies seriously. Policies and procedures are often only as good as the paper they're written on. Employers really need to be mindful of putting these policies into practice at every stage.

Victoria Ranieri: Yeah, you're totally right, Laura. For example, going back to your step one that you were just talking about, about using and learning your AI tools. When you're deciding what AI tools to use, employers want to make sure to vet their AI vendors. We're at a stage now where a lot of employers are really starting to make an investment in ai and selecting the right tools for any business [01:00:00] is key.

No matter what your industry, there are some key things that employers should understand about these potential tools before they decide to engage them. First, employers need to understand what security and data privacy protocols the vendor has in place to prevent customer data or company data from informing the broader model. In other words, maybe you don't want AI to train on your data or store your data.

Next, employers want to understand how the AI tool is trained. Employers want to make sure that the vendor has policies to train AI on reliable sources. AI really is a garbage in, garbage out model. Like you said, it can't think for itself. If you put garbage into it, you're not going to get a good product out of it.

Finally, employers should understand how the vendor is checking for and mitigating potential bias during deployment. Vicki talked earlier in the podcast about when the algorithm goes wrong, and it can go wrong, you want to know what [01:01:00] the developer is doing to prevent that when the algorithm gets in the hands of the user. Check whether the vendor has established guidelines to check for and filter out any discriminatory inputs or outputs.

Laura Corvo: Once you've selected your tools and negotiated with your vendors to make sure that they're the right tools for you, you also need to train employees on how to use them.

Victoria Ranieri: Yeah, one hundred percent. As with any new tool, employers need to not only show employees how to use AI tools, but how to use them properly and also within the parameters that the employer set through that robust AI policy that you talked about. Employers need to be really clear about the pitfalls of misuse and the consequences for misuse. Employers also need to enforce those policies and procedures and not look the other way when employees stray from them.

Like Vicki said and like we keep coming back to, AI is a simulation of human intelligence. It's not a substitute for human [01:02:00] intelligence, so employers should always make sure that humans are overseeing all of the AI tools that they use and that humans have the last check and the last word on the work product that's going out.

Okay, so now we've written our policies, we've selected our tools, and we've trained our employees. We're done. Right? We can just let the AI run now. Or are there some other things that we should be thinking about?

Victoria Fuller: I think number one, as Laura talked earlier about making sure that you know about the laws that are out there, this is a moving target. As we talked about, there are a few laws out there [and] there are more coming. Again, multi-state employers, make sure that you have an open channel with your counsel to make sure you're aware of where these laws are popping up. [Make sure] that you are in compliance with them and that you know what's coming down the pike.

Again, we talked about this earlier, you need to assess the impact of AI for potential discrimination or other [01:03:00] unlawful conduct. I'm just going to plug in here. I think this comes up almost every episode, but make sure you have insurance. If you're going to use a tool like this and you're not 100% sure about the output, you want to make sure you have an insurance policy standing behind you. Make sure that you are prepared to deal with a claim as a result of the AI tool generating a discriminatory output. If you're not prepared for that potential outcome, I would not use the tool until you can either get comfort that that's not going to happen because you've checked the results, or you've got insurance or you're able to self-insure against it.

Again, get counsel. Don't use AI. Your lawyer has gone to law school, has years of training and experience, and is not just a computer program that you know is pumping out output only based on the limited information that has been fed into it. Your council can [01:04:00] help you to draft an AI policy, can help you respond to demands or suits generated by AI. Please don't use AI as your lawyer. Chat GPT is not a lawyer.

The other thing is you don't want to feed confidential information about a claim or potential claim into AI, particularly an open AI program where you could compromise the confidentiality of that information. Just be aware that this is, as Laura said, an evolving area of the law. It's really unsettled in several different areas. As a result, we just want to be proactive and conservative in our approach so that we're not making novel law like a breach of contract claim.

Victoria Ranieri: Circling back to what you said, I want to sort of expand on some of the things you said about getting insurance because a lot of things are changing in the insurance context to try and keep up with AI as well.

You may think you have insurance for something, but you need to keep checking your policies or work with your [01:05:00] broker to check your policies. For example, in your Mobley example that you talked about where the algorithm was accused of discrimination, most employers may say, "Look, I have insurance that will cover me for discrimination," but does it cover you for discrimination by an algorithm? That is language that may be changing, coming up, because insurers are going to evaluate these risks differently, potentially. They may start accounting for that in policy language. As AI keeps developing, you want to keep making sure that the policy you have today will cover you in the event that an error is made by an algorithm and not a human.

Also, when Laura talks about AI policies, that kind of dovetails with insurance as well. A lot of underwriters are now trying to evaluate the risks that they're facing with employers who are using these AI tools, right? They're going to want to see your policies a lot of the time. They're going to start asking questions. We see this coming where they're asking questions about [01:06:00] how your business uses AI. What protections [do] you have in place to prevent the misuse of AI.

We're starting to see this, actually, with some cyber policies already. Vicki, you talked about the increased sophistication of spoofing and phishing attacks. [There's a need for] a lot of social engineering coverage. Now in the underwriting process, the cyber juror will ask the insured, "Do you have a procedure under which you call the person providing you with wire instructions

to verify that this is them on the phone and not through the computer.” Having these robust policies in place will really put you in good stead with your insurance.

It's an evolving field. Keep working with your broker and keep, as Vicki said, making sure that you're comfortable with the level of coverage that you have. [Make sure] that it's going to be able to support the work you're doing and the tools you're using.

Laura Corvo: Vicki, what else are we seeing on the horizon as we start to navigate this new AI frontier?

Victoria Fuller: Well, in the legal world, we talked about the laws that are coming out. [01:07:00] There are also judges who are issuing standing orders regarding the use of AI. Again, many of them want disclosure if you're using them. If you are, they also want certifications from either council or the party stating that they've personally checked all the citations in the final product and the output. There are maybe two dozen that I'm aware of, I think. I suspect we'll start seeing more and more of these across jurisdictions as this continues to snowball as a problem.

Those notice requirements that I just mentioned, that's not, again, unlike the laws that we're seeing starting to come out. I would expect more states to adopt laws that require notification requirements. If they're using AI as a business, they've got to disclose that to your employees, to your customers, whoever is on the receiving end of it.

The other thing is AI is affecting legal spend. We're already seeing discussions about this. Laura, can you actually talk a little bit about that [01:08:00] issue?

Laura Corvo: Yeah, I mean, I definitely think because of all the things we've talked about here today, we're going to see an increase in legal spend. There's already estimates for the legal spend being up significantly in 2026 because of, among other things, AI issues. I think employers have to realize that. Make sure, as a result of that, that they have insurance in place that could potentially cover them.

People don't want to spend money on lawyers, but spend it wisely. Bring your lawyer in on the front end of things. Get the policy set up. Do the homework so that you're not getting the big claims on the back end of things. I always tell clients, “Spend 10 minutes with me now as opposed to 10 years with me later litigating a case.” It really makes sense to get ahead of this and to take some

deliberative time to really understand how you're going to integrate AI into your business.

Victoria Fuller: Totally agree. I feel like we've covered a lot of [01:09:00] territory today.

There's also a lot more territory that we could cover, but we're really out of time. I hope that this episode has been insightful and useful for our listeners. I want to thank our listeners for joining us here today on *The Employment Law Counselor* podcast, where we talk about the risks facing employers today and discuss how better mitigation equals less litigation.

If you enjoyed this episode, please leave us a five-star review. Please tell your friends and subscribe to the podcast. For more information on this and other topics, please visit our website at White Williams, www.whiteandwilliams.com, or you could visit our blog and learn more about the firm. Until next time, bye everybody.

PLUS STAFF: Thank you for listening to this episode of *The Employment Law Counselor*. If you haven't checked out the previous episodes, make sure to give those a listen and check back in and the next few weeks for the next episode. If you have an idea for a future PLUS Podcast, you can visit the PLUS website and complete the content idea [01:10:00] form.