



Lowenstein Sandler's Insurance Recovery Podcast: Don't Take No For An Answer

Episode 109:
**Deepfakes and Insurance Policy Hot Takes:
Securing Coverage for AI Risks**

By [Lynda A. Bennett](#) and [Jeremy M. King](#)

OCTOBER 2025

-
- Lynda Bennett:** Welcome to the Lowenstein Sandler podcast series. Before we begin, please take a moment to subscribe to our podcast series at lowenstein.com/podcasts. Or find us on Amazon Music, Apple Podcasts, Audible, iHeartRadio, Spotify, Soundcloud, or YouTube. Now let's take a listen.
- Lynda Bennett:** Welcome back to *Don't Take No for an Answer*. I'm your host, Lynda Bennet, Chair of the Insurance Recovery Group here at Lowenstein. And with me is our newest addition, Jeremy King, a fellow policyholder advocate who is a seasoned insurance litigator and a staunch advocate for policyholder rights. And we're absolutely thrilled not only to have you join the firm, Jeremy, but also to join our show today.
- Jeremy King:** Thank you so much, Lynda. It's really great to be here, really great to be on the show and frankly, really great to be part of such a talented team here at Lowenstein.
- Lynda Bennett:** All right. Well, let's get to what we're going to be talking about today. We're going to focus on how AI is changing cyber risk and how insurance is evolving to address that risk. What we'll do today is we'll hit a number of the key coverage issues, the newest AI-driven threats, and what we're going to dive deeply into, I think, Jeremy, is where policy language is changing and what policyholders really need to know about now. So, to get things started, why don't we talk about the policies that matter the most for AI and cyber-related risks?
- Jeremy King:** Well, certainly. Lynda, a company's core protection usually sits in its cyber insurance policy as you'd expect, but crime coverage is also a complement for these risks. However, new AI exposures, they also can hit technology errors and omissions policies, media liability policies, frankly, even employment practices liability policies or directors and officers policies. Each of these coverages are getting new AI-specific definitions, endorsements, and exclusions across all of these lines.
- Lynda Bennett:** So, in other words, we need the whole alphabet soup of insurance coverage and we have to, as we always say on *Don't Take No for an Answer*, need to be

guided by the precise words in the insurance policy because it always matters, right?

Jeremy King: Absolutely. Reviewing the entire program for these risks is really important.

Lynda Bennett: All right. Well, everybody knows that AI is revolutionizing how companies are doing business. We're all feeling the pressure of that, but what are the new high-impact AI risks that policyholders, in particular, need to consider when they're addressing their insurance programs? This is a timely topic right now since many companies are going to be heading into their next renewal cycle.

Jeremy King: Absolutely. Look, and the bottom line here is that AI is expanding the risk, both the vectors of risks that the companies are facing and the magnitude of the potential liability that can come from those risks. Focusing on specific policy wording and renewal is really important for a company.

What we're seeing a lot of now is AI-created phishing and deepfake scams. The AI and technology's able to make these scams both harder to detect, more realistic, more convincing, and also used to manipulate audio and video as well as just email so that they can pull their tricks on HR and finance and executives inside of companies.

We're also seeing AI technologies accelerate in the hacking and hostile activity against computer systems. This makes for both faster and more sophisticated attacks that can do greater damage once they're actually inside of a company's computer system. The AI vectors can also create disclosure and risk of errors when companies are relying on AIs in particular to interface with the public, which can create an entirely new type of liability for a company they frankly haven't grappled with before.

And regulation of AI is creating regulatory and reporting concerns regarding the use of the technology. The advent of the new regulations and SEC rules means that the C-suite really needs to be concerned more than just about how AI is enhancing productivity because there could be a liability that comes from that use. And of course, AI's conduct in general may create some liabilities and there's scale risk to companies too. AI may make you more productive, but with more productivity comes more areas and more exposure to risk. So, the limits you had last year may not be sufficient for what's coming in the next few years.

Lynda Bennett: Yeah. Well, it feels like the scammers and bad actors are always two steps ahead of companies on risk management. And also, the insurance industry, as you and I both know, which is why all of the coverage grants, whether it's found in a cyber policy or in a crime policy or the others that you mentioned at the top of setting the table here, we need to keep a careful eye on the words of these policies because they are changing every single year to account for these new and emerging risks.

So now I want to pivot, Jeremy, to what are some of the new risks that are creating gray areas that are sparking coverage disputes that you see right now or that are just around the corner?

Jeremy King:

Lynda, what we see coming around the corner are right now prompt injection attack, a kind of prime example of a new gray area. These are attacks where sophisticated threat actors can craft prompts in an attempt to trick a public-facing AI into revealing what may be restricted or protected data. And all of this happens without any breach or unauthorized access to your computer systems. So, if your policy for cyber coverage triggers on something like a security breach or unauthorized access, well, you may have a coverage problem then because this AI disclosure is happening without a system intrusion of any sort.

So, you really need to look for language in these policies that addresses the data input into AI systems and treats unauthorized disclosure via that AI tool as a covered event. And this is particularly important for healthcare to financial services, but frankly, any company that may use an AI chatbot and have that chatbot with access to restricted or personal data.

Lynda Bennett:

Well, and it seems like so many companies are being encouraged very strongly to figure out ways to integrate AI into their business models in every way, shape, and form. But I think you'd agree with me that AI has really supercharged phishing and business email compromise types of scams. So, can you touch on that as well as that's certainly another huge risk area that touches frankly every industry?

Jeremy King:

Well, absolutely. AI is supercharging is a good way to put it, supercharging these attacks. It's helping to eliminate those telltale signs that make those scam emails stand out. It's getting rid of bad grammar and odd phrasing that may show up and indicate to you that an email is phony. But frankly, even worse, AI has the ability to scrub information from public data and social media so that it can make incredibly convincing, fraudulent emails that really impersonate a specific writer, contain personal details, and can imitate that writer's writing style to really enhance the scam. It makes it much easier for threat actors to tailor the fraud to a specific intended audience and makes it that much more dangerous because of that.

Some carriers are split over this kind of email intrusion. Some cover AI generated phishing and social engineering, but others can exclude it. Still more carriers will have very specific payment verification protocols that are needed as a condition to coverage, which is something you really need to look at ahead of time when you're buying coverage.

Policyholders will also find sub-limits on this type of social engineering coverage, which means that you're not covered up to what you thought of as your full limits in your cyber policy. So, you don't want to be surprised to learn when an event happens that perhaps you're not quite covered to the extent you thought you might be if you fall victim to one of these attacks.

Lynda Bennett: Yeah, who knew we would long for the days of the email from the king of Zimbabwe trying to get a check cashed here in the States. What's really shocking about these social engineering schemes and the email compromises is how long these threat actors sit on the computer system before executing on their scam. They are incredibly patient, and as you said, they're incredibly detail-oriented to follow the whole flow of who authorizes what language they use to get authorization on payments. It's really quite something. But as we said just a few minutes ago, we also know that these bad actors are working very hard to be two steps ahead of us at all times. So now we're getting better about the next scam of sitting on our email system and watching what words we use.

But let's talk about deepfakes. As we record this episode right now, the entertainment industry is up in arms about Tilly Norwood. So those of you who don't know who she is, she may put high-priced actors right out of business for work because that is an AI-generated actress, but that effort is at least transparent. What about those criminals in the cyber underworld? How are those bad actors using AI to level up from impersonating someone on an email?

Jeremy King: Well, Ms. Norwood is a good example here because she takes that step across that uncanny chasm and really presents a realistic actor. And that's not the future, I mean, that's happening right now. The AI deepfake is almost as scary from an insurance standpoint as it is from the security standpoint. I mean, there was a fairly well-publicized instance last year where an engineering firm wired \$25 million to fraudsters after an employee received a direction to disperse the funds during a video conference call where multiple attendees on that call were deepfake AI video imposters of coworkers.

Lynda Bennett: Hold on. Jeremy, I'm actually speaking to you, or am I speaking to a deepfake?

Jeremy King: This really is me, and you can tell because my mannerisms couldn't possibly be imitated by AI, but it's incredibly scary. You never can be quite sure. If your coverage hinges on network or privacy breaches or defines fraudulent instruction narrowly to focus on email or written communication, the deepfake loss, the person on the video with you that's not a real person, that could fall outside of your relevant policy language.

It's important on renewal, especially in this emerging AI era to review your definitions of your social engineering and your fraudulent instruction coverage to make sure that diverse media is covered and it's not just the transmission via email system or something like that. Because the video deepfake is unfortunately here and something we're going to need to deal with.

Lynda Bennett: Well, if we haven't scared our listeners enough yet, AI can also be used by threat actors to attack the company's computer system directly. So, does that use of AI raise a new question for cyber breach or ransom coverages for those that have dedicated cyber policies?

Jeremy King:

Now, it's certainly important to both the cyber breach and ransom coverage. On cyber breach side, it's not so much the AI, although that is a concern, it creates a much larger risk of loss and a much faster attack to the system once the AI can actually get access to it. But the actors themselves can be very concerning. We know that foreign state actors will employ AI tools and that policyholders can suffer loss from the use of those tools against their systems. Coverage may be challenged by insurers because many policies carry acts for hostile or war-like activity, which is always implicated when you have a foreign state actor that may be using any kind of tool, much less an AI tool.

For ransomware events, the use of AI can create a whole new area of exposure. Expanding AI technology has resulted in many policies now addressing extortion threats involving manipulation or corruption of the AI system itself. For instance, two months ago, one insurer introduced LLM jacking coverage, which covers losses resulting from threat actors accessing and exploiting cloud-based large language models or AIs in order to create erroneous outputs and corrupt the AI model. Essentially, they're saying, "We're going to hold your AI hostage or pay us the ransom."

They also have similar schemes involving LLM jailbreaking, which is removing or bypassing the safety restrictions that allow the model to operate. These situations involve both extortion losses but also business interruption losses for companies that rely on these models and could result in direct loss to the business itself. So, your cyber ransom and cyber extortion coverage should be broad enough to cover this kind of LLM jacking event, which may impact your business's bottom line at the end of the day.

Lynda Bennett:

Up to this point, we've been largely talking about those direct company threats. Let's turn to liability for a second. When a company's AI gives bad advice or it hallucinates, the legal industry certainly are learning all about this as we see near daily reports of lawyers getting sanctioned for filing briefs with fake citations. And some courts have even fallen prey to this and have issued opinions that they've had to take down because it contains bad case citations. So, what kind of coverage might there be available to respond when a third party has been damaged by a company's negligent or irresponsible use of AI?

Jeremy King:

Typically, that would fall within a technology errors and omissions policy or a professional liability policy, especially when you're dealing with a profession such as the practice of law. And these cover claims made against the company arising out of reliance or the use of that advice. But this potential liability does raise another gray area about just the definitions of technology services or professional services and whether AI systems are explicitly included in providing those services.

So, if the AI acts like an agent for your company to the public, you need to make sure that your insurance treats its outputs in the same way that it would treat the output of an insured employee or a vendor, not just a product or software, which AI could be characterized as. The insuring agreements covering the act of the

employee really need to also include that AI is telling people and what your clients are relying on your AI to tell them.

Lynda Bennett: All right. Well, there's certainly more potential liability risks, but in the interest of time, Jeremy, let's move on over to regulation of AI. That seems to be ramping up too. What do policyholders need to be thinking about in terms of insurance coverage when the regulators come knocking?

Jeremy King: The laws impacting AI usage are proliferating around the country. California has passed Defending Democracy from Deepfake Deception Act and an AI transparency act. There are privacy protection agencies proposing rules on cybersecurity and risk assessments. Many other states are following suit and adding their own regulatory regime. But the regulatory coverage that you'll find in your AI policy, that typically ties to security or privacy events, breaches that you may have an obligation to report.

In your policies, you should look for updated definitions of privacy law or regulatory investigations to include these AI regulatory regimes and decision-making rules that are being implemented around the use of AI systems. Some policies will only cover the regulatory fines or investigations for AI misuse, but they may not include certain penalties or AI-specific statutes that a company may have to perform under. You want to bridge any potential gap here with respect to the regulations through D&O or E&O coverages.

If your D&O insurance contains broad AI exclusions, you might be able to create a solution via your cyber endorsements. One cyber insurer, for instance, offers an SEC reporting cost endorsement that will cover the cost of complying with SEC cybersecurity reporting rules. The important part is to make sure that these two areas of coverage work together during your underwriting and your renewals rather than operating independently.

We've also seen D&O coverage potentially come into play with new claims of AI washing. These are statements regarding the implementation or use of AI that are potentially misleading to investors. Current SEC rules focus a lot on the disclosure of cybersecurity incidents, your risk management strategy, and your corporate governance. But the proliferation of claims dealing with AI misstatements related to the business could also start to trigger D&O policies. And you want to make sure you don't have broad exclusions in those D&O policies that would prevent those claims from being defended or covered.

Lynda Bennett: Well, that's great, Jeremy. And I think you've done a masterful job in this episode setting the table for us on what are the new and emerging risks that are out there as well as the insurance policies that are currently available to cover these types of risks.

As I mentioned at the top, and as we always discuss on *Don't Take No for an Answer*, the precise words of your insurance policy matter. And I'd love to have you come on back and continue this very interesting conversation where we can

talk about, as I mentioned at the top of the episode, folks are going to see some pretty significant changes coming up on their renewals. So why don't you come on back another time and we can talk about the changes to the underwriting process and what policyholders need to do to ready themselves for that? But thanks for coming on today and sharing your very extensive knowledge on this very interesting topic.

Jeremy King:

Lynda, I was very, very happy to be here. And I'd be happy to come back, either myself or perhaps my AI deepfake who can take my place.

Lynda Bennett:

Thank you for listening to today's episode. Please subscribe to our podcast series at lowenstein.com/podcast or find us on Amazon Music, Apple Podcasts, Audible, iHeartRadio, Spotify, Soundcloud or YouTube. Lowenstein Sandler Podcast series is presented by Lowenstein Sandler and cannot be copied or rebroadcast without consent. The information provided is intended for a general audience and is not legal advice or a substitute for the advice of counsel. Prior results do not guarantee a similar outcome. Content reflects the personal views and opinions of the participants. No attorney-client relationship is being created by this podcast and all rights are reserved.