

NATIONAL COUNCIL OF INSURANCE LEGISLATORS
FINANCIAL SERVICES & MULTI-LINES ISSUES COMMITTEE
2026 NCOIL SUMMER MEETING – BOSTON, MASSACHUSETTS
JULY 16, 2026
DRAFT MINUTES

The National Council of Insurance Legislators (NCOIL) Financial Services & Multi-Lines Issues Committee met at The Westin Copley Place in Boston, Massachusetts on Thursday, July 16, 2026 at 3:45 p.m.

New York Assemblyman Jarett Gandolfo, Chair of the Committee, presided.

Other members of the Committee present were:

Sen. Justin Boyd (AR)	Asw. Pam Hunter (NY)
Sen. Tim Grayson (CA)	Asm. David Weprin (NY)
Rep. Michael Meredith (KY)	Rep. Brian Lampton (OH)
Rep. Michael Sarge Pollock (KY)	Sen. George Lang (OH)
Rep. Edmond Jordan (LA)	Rep. Ellyn Hefner (OK)
Rep. Shaun Mena (LA)	Sen. Mark Mann (OK)
Rep. Brenda Carter (MI)	Rep. Tom Oliverson, MD (TX)
Sen. Lana Theis (MI)	Rep. Trey Wharton (TX)
Sen. Paul Utke (MN)	Del. Walter Hall (WV)
Asm. Erik Dilan (NY)	

Other legislators present were:

Rep. Stephen Meskers (CT)	Rep. Robert Foley (ME)
Rep. Kerry Wood (CT)	Del. Pamela Guzzone (MD)
Rep. Elijah Pierick (HI)	Rep. David LeBoeuf (MA)
Rep. Rita Mayfield (IL)	Sen. Mark Huizenga (MI)
Sen. Willie Preston (IL)	Sen. William Gannon (NH)
Rep. Peggy Mayfield (IN)	Sen. Tim McGough (NH)
Rep. Dennis Bramburg (LA)	Rep. Kevin Norwood (OK)
Rep. Aimee Freeman (LA)	Rep. William Slater (TN)
Rep. John Illg (LA)	Sen. Cale Case (WY)
Sen. Kirk Talbot (LA)	Rep. Pepper Ottman (WY)
Rep. Michelle Boyer (ME)	
Rep. Rolf Olsen (ME)	

Also in attendance were:

Will Melofchik, NCOIL CEO
Christa Rapoport, NCOIL General Counsel
Pat Gilbert, Director of Policy, Administration & Member Services, NCOIL Support Services, LLC

QUORUM

Upon a Motion made by Asm. Erik Dilan (NY) and seconded by Rep. Brian Lampton (OH), the Committee voted without objection by way of a voice vote to waive the quorum requirement.

MINUTES

Upon a Motion made by Asw. Pam Hunter (NY), NCOIL Immediate Past President, and seconded by Rep. Tom Oliverson, M.D. (TX), the Committee voted without objection by way of a voice vote to adopt the minutes of the Committee's April 19, 2026 meeting.

PRESENTATION ON INSURANCE AND ARTIFICIAL INTELLIGENCE – LEVEL SETTING AND DISCUSSING THE ROAD AHEAD

Asm. Gandolfo stated that first up on our agenda today is a presentation on Insurance and Artificial Intelligence (AI) - level setting and discussing the road ahead. I think it's important to continue to have discussions about AI throughout the year as there's still so much to discuss and learn on the topic and today, we have a very interesting presentation on insurance and AI on our agenda. There are still so many unknowns surrounding AI and insurance, and as the technology continues to advance so quickly, it can clearly provide a wide array of benefits to both insurers and consumers but some are concerned around things like unlawful and unfair discrimination, underwriting declinations or claim denials without a human in the mix. Throughout the country, bills continue to be introduced in state legislatures at a rapid pace concerning insurance and AI. So, in light of all this activity, this is a great opportunity to level set some things surrounding insurers' use of AI so that we as policymakers better know and understand how the technology is being used.

Steve Armstrong, SVP and Chief Actuary at Allstate, thanked the Committee and stated that I truly appreciate this time to spend with you on a topic that is moving quickly, creating real opportunity and raising fair questions for policymakers, regulators, companies and consumers alike. I want to start by saying what this session is not going to be. First and foremost, there will be no actuarial formulas or equations. This is not going to be a technical lecture. Number two, this is not going to be a tour through every algorithm, every model or every new AI tool that's out there in the marketplace. And this lecture is not going to be a discussion about where we treat AI like one giant mysterious black box. Instead, I aim going to try to keep this conversation practical. Reason being because for consumers, AI is not an abstract technology debate. AI shows up in moments that matter every day in their lives. In insurance, AI may show up when someone is shopping for insurance. It may show up when they're trying to understand coverage. It may show up when they're trying to get a quote. It may also show up when a claim is routed or reviewed, and it can show up when a company might try to explain a price change, a coverage issue, or a claims decision. So, today my preference is to ground all of this discussion in one person. I want to call her "Maria". Maria is going to be the focal point of our conversation today. Maria has just moved. Maria is shopping for auto insurance and homeowners insurance. She is busy like everybody is. Maria is trying to make a good decision, but she may not know all the insurance terms that exist out there. She wants a fair price. Maria wants coverage she understands, and if something does not make sense to her, she wants someone to explain it in plain language. And that's going to be our anchor today. As we go through the discussion, I'm going to ask you all to keep Maria in mind as a reminder that insurance decisions affect real people at real moments.

And I want to also ground us in three questions to hold on to. I'm going to repeat these questions throughout the presentation. The first question that I want everybody to keep in mind is, what is AI doing? Second, does it affect Maria? And third, can the company explain the AI and stand behind it? I truly believe if we can answer these three questions, we can have meaningful policy conversations. You all, contrary to what I like to do, do not need to become data scientists to ask strong oversight questions. Sometimes the strongest questions are often

the plainest ones. For example, what is the tool doing? Does it affect the consumer? And who is responsible? And that's going to be our frame for today. We're going to start with the consumer and talk a little bit more about Maria. I'm going to attempt to demystify AI a little bit. We're going to talk about different levels of consumer impact. We're going to talk about both the benefits and the concerns of AI to a consumer. They're important because they're both real. Then we're going to talk about accountability as well as trust. And then we'll end with what I hope is a nice practical watch list that you can use in legislative, regulatory, and/or oversight conversations.

So, where should every conversation with AI begin? It should start with the consumer. Maria does not walk asking whether or not a tool is machine learning or generative AI or predictive analytics or agentic AI, she obviously does not care about the label. What Maria does care about is the outcome. She cares whether that coverage makes sense. She cares whether that price is fair. She cares whether the company has used information appropriately, and she cares whether someone can explain the answer if she has a question. So, that brings us back to that first question. I want everybody to think through what is the AI doing? Maybe AI is helping Maria compare coverage options. Maybe, it's translating that goofy insurance language that sometimes doesn't mean a lot into plain language that Maria can understand. Maybe it's also helping a company process her information more quickly. Maybe it's supporting underwriting or pricing analysis. Maybe the AI is helping route a claim. Or maybe it's helping support a decision that can materially affect Maria. The point here is all those maybes are very different things. So, what is AI doing is question number one.

Question number two, does it affect Maria? Does the AI touch her price? Does it touch her eligibility? Does it touch her coverage? Does it affect her underwriting? Does it affect a claim? Does it affect her renewal offer? If the answer is no to many or all of these questions, then the oversight question might actually be on the light side. However, if the answer is yes, then the accountability question becomes much more important. And, then that brings us to our third question. Can the company explain it and stand behind it? This isn't just about producing a big technical report but telling everybody what the AI is doing. It is about whether there is a clear explanation, a responsible owner, and a path for review and correction. Maria should not be trapped in a maze where every answer is the system said so. Because that is obviously not accountability. A good accountability framework should make sure that when AI affects a consumer's outcome, there is still a company and real people within that company responsible for the result.

So this slide is an attempt for some potential audience participation. I'd like to have us all put ourselves in Maria's shoes for a moment. Maria is in the process of shopping for insurance. She takes out her cell phone, goes to her favorite AI of choice, and she types in, "Hey, I just moved and I need auto insurance and homeowners insurance. I do not know much about insurance. Can you help me understand what coverage I might need? What affects my premium? And what questions I should ask before I buy?" This is a very realistic future. Frankly, it's already becoming a realistic present. People are looking to shop for auto insurance with their phones more often. And so, consumers are obviously going and using AI before they ever go to a carrier website, before they ever talk to an agent, before they ever visit a comparison site and before they actually get a quote.

The poll up here is a question if you were Maria in this situation. Which of options A, B, C, D or E and it could be all, none or many. There's not one right answer, but what do you think is which AI supported use case would be most helpful to Maria as she's beginning shopping for insurance? So, A: AI explains the coverage in plain language. B: AI suggests questions to ask before buying. C: AI helps identify maybe a coverage gap that she didn't know about. D: AI

explains what affects her premium? E: AI tells her what company and price she'll pay? So, arguably all of these sound helpful. I think all of these would help me if I was shopping for insurance, but arguably each one of these five do not carry the same consumer impact. So revisiting number A, if AI explains coverage options in plain language, that is useful because it helps Maria become a more informed shopper. Going to option B, if AI suggests questions to ask, that's also helpful because it will make Maria feel more prepared. Going to C, if AI helps identify coverage gaps, that can be very valuable for Maria, especially since she's someone who just moved and may not know the local risks or coverage differences in the state. Now for D, if AI helps explain her premium that can help kind of demystify insurance and have her trust it more. And with E, if it recommends a company and an estimated price, that might be where I'd ask you to pause and go, how helpful is that to Maria? In theory, it should be. But now AI has started to move away from education and a little bit more into recommendation. And that recommendation may sound really confident, even though it may not necessarily reflect an actual filed rate, current underwriting eligibility, company's coverage details, or the most current information. And therein lies the caution here. As AI moves from education to recommendation, the need for explanation, accuracy verification and accountability increases. A simple thing to remember is that AI can be a map, but it should not be necessarily the destination. If Maria uses the AI to understand that landscape it can certainly help her. But if AI starts telling her where to buy, what it will cost, and what she qualifies for, then the questions become a little bit more serious. Because what information is AI using in this case? Is that answer even accurate? Can Maria verify it? And who's accountable if the answer is wrong?

Moving on, let's talk a little bit more about Maria. She just started shopping for insurance and she asked AI for some help. Now Maria she can shop for insurance. Maria gets a quote, she chooses coverage, she receives a policy, she eventually makes a claim and throughout she might need explanations for questions that she otherwise has. The point of this is that these are all very different moments in an insurance journey and AI can and will show up very differently in each one of these. So, I'd like to overly simplify and put these into two different categories. On the left-hand side you have a category of AI that supports Maria directly as a consumer because it might help her get started. It might help her compare those coverage options. It may help her understand what uninsured motorist coverage is. It may help her get that quote. It may help reroute her request to the right place, and these uses can make the experience for Maria faster, clearer, and obviously less frustrating. On the right side though, the category is it supports company decisions. This might include helping the insurance company with quote and underwriting information, AI might support pricing and eligibility decisions. It might influence coverage recommendations. It might help triage a claim. It might help support claim handling or even provide a claim payment, and it could help obviously generate explanations.

This second category here on the right is the one that deserves special attention because it's closer to decisions that affect Maria's outcome. One of the mistakes sometimes I hear is that in AI and insurance is that it's just "AI insurance" as if that's all just one big thing. And I'm here to argue that I don't think it is one thing. AI is a set of tools used in different places for different purposes that require different levels, or that have rather different levels of consumer impact. So, the practical question to then ask again is, where is AI in the journey? Is it helping Maria understand? Is it helping the company do processes? Is it supporting a decision, and is it shaping an outcome? These are very important questions, which I'll come back to in a few more slides. Again, AI is not necessarily one label, but sometimes we see these three terms interchangeably used. So, I thought I'd maybe kind of talk about each and try to offer a policy lens in consideration for each.

Starting on the left, machine learning and predictive analytics are definitely the most established of all the AI on this page, because insurers and actuaries like myself have been using data, mathematical models and professional judgment for a very long time, decades actually. The tools have changed, but the underlying discipline of analyzing risk, testing assumptions, validating models, and applying judgment is nowhere new. Predictive tools rather may help identify patterns in data that can't be seen. And in insurance, we kind of boil that down to saying it's supporting risk assessment. So, the policy lens when looking at AI and the lens of machine learning and predictive analytics is really how well is the AI assessing that risk? And maybe even more importantly, how fair is that assessment? Moving more into the middle, generative AI or Gen AI as it's often called, is very different. Obviously, it can create or summarize text. It can explain things in plain language. It can turn complex information into something easy to understand. And here I would argue that the policy lens is how accurate is it, and what review process or review exists, because a generative AI answer can sound very polished, even when it can be completely wrong or incomplete. And then moving to the far right on this page, we have agentic AI and this is the one that's emerging and this is where tools may eventually take steps across systems, route work, trigger different workflows or act independently as well. The policy lens here is authority and responsibility review because if a tool can take action, we need to be clear about what authority it has and who is responsible ultimately for that decision. Bringing the discussion back again to ask, what job is the AI doing? It could be finding patterns. It could be summarizing. It could be explaining. It could be recommending. It could be routing. It can be supporting a decision or it can be taking an action. The label doesn't matter; the role and the impact of the AI are what's important.

And then bringing that back to our customer Maria one more time, what matters is not whether the tool is called something. What matters to her is that the tool affects her quote, her coverage, her claim and her ability to get a clear answer. I would argue that this slide and what's coming up are maybe, in my opinion, the two of the most more important slides. Simply said, slide seven says different impact, different accountability. There is a spectrum of impact. On the low on the left side, on what we're calling lower impact end-AI can be used potentially for internal productivity. Could be drafting a memo, sorting documents, helping with some coding, summarizing information for employees. These uses still do need appropriate controls, but they are generally further away from a direct consumer outcome. As we move over to the right, we have things kind of in the category of consumer guidance. AI might help explain coverage and summarize those options. This can be very helpful as we discussed, but accuracy does matter. If Maria gets an explanation that is unclear or wrong, that can make for a very bad decision. As we move further to the right, we have a category called decision support. AI might support pricing decisions, eligibility decisions, underwriting and claims review. Now, we start to see that the consumer impact grows and the potential for negative impact becomes much higher.

And finally, on the right, I have what I'm terming "consequential outcomes." These can be non-renewals, claim payments, material premium changes or other decisions that meaningfully affect the consumer. And the rule of thumb here is very simple; scrutiny should rise as AI gets closer to material consumer outcomes. And that's the practical principle here every day, and it helps avoid treating every AI function like it carries the same risk. And it avoids pretending that high impact AI uses are just ordinary automation. The spectrum's important to determine where that scrutiny threshold plays out over time. And then bringing it back again to Maria, AI that explains a coverage term is different from AI that affects her when she might be looking for coverage, or what she pays in premium, or what happens when she files a claim. The more AI moves the outcome, the more the company needs to be able to explain it and stand behind it. That does not mean, though, that every step needs to be manually approved by a person. But it does

mean that the controls, the testing, the documentation, the explanation and the escalation should match the level of impact on the consumer.

The next two slides address how AI can help consumers, and how AI can harm consumers. Two sides of the same coin. So how can AI help Maria? If used well, AI can make Maria's experience faster, clearer and more consistent. When Maria is shopping, she can compare options side by side. It can help explain coverage in plain language. It can speed intake and review, so she is not stuck repeating the same information across multiple channels. And once she has that coverage, arguably AI can help route a claim more efficiently. It can help explain what changed and why, and it can support smoother service because in insurance we all know clarity matters, speed matters, and consistency matters. We also know many consumers find insurance confusing. They may not understand what a deductible or a limit or endorsements or exclusions are, and AI can help translate that complexity into plain language, and that's no small benefit. If Maria as a consumer can better understand what she's buying and the advantage to her, she's more empowered. If she can get a faster answer, she's less frustrated. If a claim or service request gets routed more smoothly that can improve her experience at a very stressful moment. In short, Maria should benefit from better tools and technology. However, AI can go wrong. The same tools that can help her, can also disadvantage her if they're not governed well, and there are several ways that this can happen. The first one could be wrong data. If the information going in is wrong, the answer coming out is very likely to be wrong. That is true, whether we're talking about AI or any model. Quite practically, bad input usually leads to bad outcomes.

The next slide addresses unfair data patterns. A model might find a pattern in the data, but the company still has to ask whether or not that pattern is appropriate, fair and legally permissible. A pattern is not automatically a good reason just because a model found it. Third, poor explanations. An answer may be technically complex, but consumers still need plain language explanation of all the main reasons. If Maria cannot understand why something happened, then trust begins to break down. Fourth, no responsible review. If something goes wrong, Maria needs a path to a real answer, not a loop, not a dead end, not the system said so, a true path. And fifth, and maybe most important is overreliance on AI. This is a real risk. The tool may be treated like the final answer instead of support for human judgment, document and review and controlled decision making. The strip at the very bottom of the slide shows what things can go off track. Because data's going to feed a pattern, a pattern can support a decision, a decision's going to need an explanation, an explanation therefore may need a review and a review may need a correction. A weak link anywhere in that chain can definitely harm Maria. If the AI affects the data, the pattern, the decision, the explanation or the review process then accountability truly needs to follow.

I'm going to now talk about accountability. Accountability is the principle that people remain responsible. When AI begins to affect Maria, real people remain responsible for the decisions, the explanations, the documentation and the review. That does not again mean a person must manually approve every AI supported action because that's not realistic, and it may not even be necessary for some of those lower impact cases. But it does mean that the company must be able to stand behind that process and that outcome. And so, again, looking at the slide here for the lower-impacted AI, monitoring and quality checks may be sufficient for accountability. But for medium-impact AI, as you may expect, which could include document review, exception handling and escalation, more controls need to be put in place for accountability. On the far right for high-impact AI, especially where AI affects price eligibility claim payments renewals or other meaningful consumer outcomes, way stronger human review and documentation are appropriate. So, the four things that I believe when you put together create an accountability package are a clear explanation, documentation, a responsible owner, and a path to review for

correction. And if companies honestly want to keep using increasingly powerful tools in consequential areas then they need to build the accountability muscle that earns confidence. For Maria again, the question is very simple. If AI affects my price, my coverage, my claim, or my renewal, can that company tell me what happened, who owns it, and how I can get it reviewed if something goes wrong? That's not too much to ask. That's what she's owed.

Moving on, I want to talk a little bit about transparency. We talked earlier that accountability is about whether the company owns the outcome. Transparency on the other hand is about whether others can understand enough to hold that company accountable. And so, the important thing here is to not say a person reviewed it because that's not enough for transparency. It may be part of the answer, but definitely not the full answer. The company should be able to explain what happened, who owns that result and what review occurred and how can a consumer seek correction. This does not mean that every consumer needs a technical explanation of the model. Most consumers don't want that. What she needs is a clear answer to these questions: What role did AI play? Why did it matter for her? Who was accountable for the results, and how can questions or corrections be handled? These are the four prompts on the slide: what, why, who, and how. Again, it goes back to what did AI do? Why did it matter? Who owns the results and how can it be reviewed or corrected. That is transferable and usable transparency. And it's not transparency for show. It's transparency that helps a consumer, a regulator, or an oversight body understand the issue and act on it. And one more very important point is that transparency also has to respect privacy. The goal is not to expose any private data unnecessarily. The goal is to identify the type of information used, explain why it was relevant, and confirm that it was handled appropriately. So, transparency is not for just dumping information; it's giving a clear, useful answer to the consumer.

Moving on, building on the framework, it's important to know that there are existing frameworks because then we might not necessarily have to start from zero. There are already several accountability structures in play. First, one thing I hold dear and near to my heart is actuarial professionalism. Credentialed actuaries are required to work under professional standards. These standards address things like risk classification, modeling, validation assumptions and professional judgment. That matters because actuaries have long worked with data and models that affect insurance operations and outcomes. Second, maybe the most important thing on here: company governance. Companies need internal controls including documentation, record keeping, model testing, validation, risk assessments, privacy controls, cybersecurity controls, and ongoing monitoring. These are not optional extras when a tool or technology affects consumers, they are part of a responsible operation. And third, state insurance oversight, state insurance regulators already have important authority through rate and form review, market conduct exams, and consumer protection tools. So, the opportunity here is to potentially build on existing oversight where it fits while also asking whether AI creates new questions that need clearer answers.

Moving on, what builds trust in this world of AI? Trust is obviously not assumed; trust needs to be earned. And when AI supports an insurance decision, I think there are four practical signals of trust for us to consider. First is a plain language explanation. Not every technical detail, not a 1,000 page model document, but a clear explanation of the main reasons that matter to a consumer. Second, clear accountability by the insurer builds trust. The company cannot point to the tool as if the tool is the accountable party, the company is accountable. Third, testing and review by qualified professionals, including credentialed actuaries where any actuary or pricing work is involved. If a tool supports pricing, underwriting, risk class or other insurance decisions, that review needs to be done by people who are qualified who understand that work. And fourth, a path for review or correction. This may be the most human signal of them all, because

consumers know that mistakes can happen. What they want to know is whether the company will listen, review and correct when appropriate. I might just want to add that none of these signals is enough by itself because an explanation without accountability is weak. Accountability without testing is incomplete. Testing without correction paths may not help a person like Maria, and a correction path without a clear explanation leaves people frustrated. But together they definitely build trust. So, for Maria again, trust does not come from the company saying do not worry we use AI responsibly. It comes from the company being able to answer here's what happened, here's why it mattered, here's who owns it, here's how it was tested, here's how you can ask for a review, and that's how I would argue trust is earned.

Moving on is what I am terming a practical legislator watch list. These are the questions that I hope can help in hearings, in oversight discussions, in bill drafting, and in conversations with regulators and companies. First again the role: What is AI doing? Do not stop at the label. Ask what the tool is performing. Number two, use: How is AI used? Is it internal? Is it consumer-facing? Is it advisory or is it part of a decision process? Third is impact. This is the hinge question. If the tool touches price, eligibility, coverage, claims, renewals or explanations, the stakes are arguably that much higher. Fourth again is accountability. Who is accountable and how can issues be corrected. If no one can answer that, that's a red flag. Fifth is an explanation. Can consumers get a clear explanation, not a technical dump. Sixth, testing - is testing and review matched to the impact. The higher the impact, the stronger the testing documentation review and monitoring should be. This watch list then supports a balanced policy focus across three different dimensions. Risk-based governance, consumer impact thresholds, and regulator access and documentation. This avoids treating every AI use as equally risky. It avoids ignoring high-impact uses, and most importantly, it keeps the consumer at the center. For Maria, this is simple: if it affects her, someone should be able to explain it; if it materially affects her, someone should own it. And if it's wrong, there should be a path to review or correction. This slide is something that you may want to be able to refer to later if helpful. In the actuarial profession, the American Academy of Actuaries, the Casualty Actuarial Society, and the Society of Actuaries are consistently publishing and researching AI and so that might be a nice resource as you're discussing this. Again, and in conclusion, the three questions I wanted to leave with you is the Maria test. If Maria's affected, can the company answer these three questions: What did AI do? How did it affect her outcome? And can the company explain it and stand behind it? That's the test because consumers deserve a clear answer and a responsible owner.

Asm. Gandolfo asked how has AI improved underwriting accuracy or claim outcomes? Are there any data points that you can provide or just a general thought on it?

Mr. Armstrong stated in terms of pricing or data accuracy, AI helps detect patterns nicely so it can detect data anomalies in ways that we might not be able to see today or humans might be missing. So it can help in terms of verifying that information's coming in and that it looks appropriate or if it doesn't, it could flag something that might need extra review. So, I think it helps with data accuracy and some of the verification opportunities there. It's also helping us surface things like in that machine learning and predictive modeling world that I was explaining - AI is surfacing risk assessments that we may have never seen before that might be appropriate but need to be tested again obviously for bias and legal permissibility. So, it's opening the doors for better risk assessment, which arguably is better for the consumers because it allows us to price them more appropriately.

On the claim side, where I've seen AI most effective at this point in time is at that crucial moment when a consumer calls in and says, "I have a claim." In many cases today, you might call the insurance company and what will happen is that you might get a claim adjuster on the

phone with you right away. That claim adjuster is supposed to treat everybody's claim through a certain protocol, the same way with the same dignity, respect, etc. But not all humans are the same. We want to have everybody have a very consistent start to their claims process. And so where I've seen it come in is helping everybody have a very consistent way of being able to say here's what happened, and here's how I reported my claim. The AI will ingest that, it'll transcribe it accurately, it'll confirm things appropriately and get everybody started on the same foot from a consistency standpoint. If a claim happens, it'll assign a claims adjuster at the same time and then what it can do is it can help you as a consumer who just had a claim stay very well informed throughout the path of the claim versus waiting for a claim adjuster to call every week and do a follow-up for you. So, it helps with consistent claims handling. That is claims intake at the very beginning and then keeps consumers more well informed throughout the claims process as well so that it can make sure that there's never a question that you have and you know everything that's going on every moment of time.

Rep. Ellyn Hefner (OK) stated you just described a process for claims. How do you disclose that to a customer or a client? Is that disclosed and what does your company do to disclose that? Mr. Armstrong stated, yes, at least at Allstate, our goal is to make sure that we know the consumer knows they're interacting with AI. So, the on front would be a disclosure like, "Hey I'm Steve. I'm your digital assistant here to help intake your claim." And so the consumer theoretically knows right away that I'm not talking to a human, but I'm talking to some AI and maybe the consumer likes it or doesn't. If they don't, they can opt out of it and reach a real person. But we try to inform them right up front that they're interacting with something artificial so that they can be aware and that may lead to different reactions and consequences.

Rep. Hefner asked Mr. Armstrong if he used AI for his presentation? Mr. Armstrong replied that he did. It helped me prep. I don't work with legislators on a day-to-day basis so the first thing I asked AI was what are the major concerns that legislators have concerning AI? And while AI can be wrong, it gave me a very good starting point to go off of.

Asm. Gandolfo stated as a broader industry trend, do you anticipate the use of AI eventually displacing some workers in the insurance industry? I know some of the items in the slide deck talked about document review and that kind of clerical work. Is there any thought on what it looks like five years from now as AI gets a little more ubiquitous in daily operations? Mr. Armstrong stated yes, but while we are cognizant that certain tasks are likely to be changed by AI in the future and that could be a part of someone's job or maybe someone's full time job, we want to think about upscaling them to do jobs in an AI world so that they don't have to lose their job. They just have to start learning new skills and do different things. And so, I think the answer will be yes, but if we do our jobs well, they will have different and equally meaningful jobs that they didn't have before. So, they might lose something but hopefully gain something in return.

Rep. David LeBoeuf (MA) stated I have a question regarding the experience with the AI agent. What guardrails do you have in place about making sure that the input into that or that interaction with the AI agent is accurately collected in the way that the consumer is actually putting it in, and the reason why I am raising this is I know there is a case in another state where a police department was using AI to enhance their police reporting, and it actually had ruined cases because of the way the AI did it. It didn't accurately capture what was said or what was visible because it polished it and it tanked a couple of cases. I know there's a couple of other states that have legislation around medical note dictating, making sure that the technology doesn't basically upcharge so that way the medical notes or something similar are coded in a way to bill the insurance companies more money. So, what guardrails are in place, whether it's

your company or the industry, currently to make sure that the actual information is not polished in a particular way to present a different picture than the consumer is presenting?

Mr. Armstrong stated I hope I understood your question well enough that this answer works. If it doesn't, I'll be happy to talk more offline with you. But I think a lot of it goes into rigorous testing protocols, quite honestly. So, it's impossible to think of every element that could possibly be out there, but as you're trying to test an agent to be able to do and perform certain tasks and take inputs and then create outputs, you have to think through all the different ways that it can be input and test whether it's the accurate output. And so it's consistent testing to make sure that it's doing what it's supposed to be doing and not doing something like polishing or potentially hallucinating or anything to that extent. And it just requires a really robust testing protocol up front and then monitoring on the back end to make sure that if it does start to slip or does start to do something that it shouldn't, that it can be caught quickly and remediated back into compliance. So, I think the answer that comes to mind for me is just rigorous testing and monitoring thereafter after the fact. So that's maybe not exactly a guardrail, but it's a requirement I think operationally to make it work. But I'm looking at you and sensing that maybe I didn't understand your question completely.

Rep. LeBoeuf stated you did. When you have a call with a human, you can go back and listen to the call in case there was some error that was made. With AI, you can't necessarily always go back and trace it unless it's designed in a particular way that allows you to. Mr. Armstrong stated, yes, that's fair. I think there are steps with a lot of AI that you have to obviously ensure that it will document what it did and why it did what it did so you can have an audit trail going backwards. So if something is wrong, you can go back and go, "Oh, this is why it went wrong because this was unclear to the agent when we programmed it. We told it to do this and we thought it was doing this, but it just didn't do it." So now you can see that from the audit and kind of correct it. So, there are some of those audits, catching things as quickly as real time as possible so consumers aren't harmed which is obviously the most important part. But there should be ways that companies can go back and audit and not just say, "It made a decision. I don't know why it did that." If you ever hear that then that's not responsible use of AI.

PRIVATE CREDIT AND INSURANCE 101

Asm. Gandolfo stated next up on our agenda is a 101 presentation on private credit and insurance. And before we begin, I'd like to note that this next presentation is closed to the press. So, if there are any press in here, which I don't believe there are, please leave now. Now, private credit is a topic that has gained a lot of attention recently as there are some concerns about insurers who hold private credit having an illiquid investment and therefore, it's a hard to sell investment that might negatively impact paying claims. But similar to what we just heard about with AI and insurance, there are a lot of unknowns and misconceptions surrounding private credit and insurance. So, this is a good opportunity today to provide some baseline understanding of some of these things.

Bridget Hagan, Managing Director & Global Head of Insurance Regulatory at Blackstone Credit & Insurance stated we really appreciate the opportunity to be here and we are going to try to keep the formal part of this presentation as brief as possible to allow as much time for questions as we can. Please don't hesitate to raise your hand as we're speaking. Blackstone is the world's largest alternative asset manager. I'll take a second just to talk about our presence in the insurance market and why it might be of interest to you all as state legislators that are focused on insurance policy. Of our \$1.3 trillion in global assets under management, about \$280 billion of that is investing assets for insurance companies. Blackstone does not own a majority stake in

any insurance company on balance sheet. We have in our insurance asset management business and what we call an open architecture model because we would like to manage assets for as many insurance companies as would like to do business with us around the world.

Today, that's about 40 or so insurance companies that have independent contractual relationships with us to manage one or more categories of assets, many of which are private credit assets. So, I am very happy to answer more questions about our place in the market but just to sum all of that up, we are the largest arm's length, third party manager of private credit assets for insurance companies and particularly in the U.S. So, with that background, before I turn it over to my colleague and we can get into some of the specific assets and specific questions, we just wanted to spend a moment on the history of how did we get here and why is this happening in insurance investing.

As mentioned, this is in the news and I'm sure you're talking to folks about it in terms of other policymakers and insurance companies, and you're I'm sure very familiar with the low for long interest rate environment that did create some pressure on insurance company investing. That's very true and very much part of this picture but one of the phenomena that we think is a little bit under discussed in this conversation is that for much of the history of insurance company investing in the U.S, insurance companies had a high percentage of assets on their balance sheet that were relatively illiquid, specifically corporate bonds. Corporate bonds essentially didn't used to be nearly as liquid as they are today, and assets that are not liquid in general feature what we call an illiquidity premium, which is just a little bit of extra yield that you have to pay investors for the fact that they're not widely tradable on an open public market. That illiquidity premium is about 1.5 to 2 percentage points of yield relative to public assets. So, hearkening back to the corporate bond phenomenon, when corporate bonds became more liquid, among other reasons technology advanced and you had electronic trading platforms, the yields on corporate bonds decreased because you weren't getting that illiquidity premium anymore. And so insurance companies really needed to find another source of yield to capture that illiquidity premium and I know you all know this as legislators that deal with insurance every day, but relative to banks, for example, that have overnight deposits, insurance companies and life and annuity companies in particular are very good holders of illiquid assets. They're not great holders of credit risk, as you know, but their liabilities, the long duration liabilities, match up well with assets that feature that illiquidity premium. We just wanted to pause for one second and explain that this decrease in the yield on corporate bonds because they became more liquid is part of why private credit is growing as a percentage of insurance company balance sheets.

So, with that background, I will just briefly talk about how we at Blackstone define private credit and I think our definition is in line with a common definition in the financial markets. You all know obviously that there is no single agreed upon global definition of private credit so it's probably worth taking a second to talk about how we view what is private credit. The most important component of private credit is that you are taking the borrower directly to the lender and matching them up directly. The distinguishing characteristic of that is that you're somewhat disintermediating the banks from that relationship. We partner with banks all the time. Banks are obviously hugely important to lending and to the overall financial markets, but there are some circumstances where it is appropriate to bring the borrower directly to the lender. What that does is cut out bank fees and so that 1.5 to 2 percentage point difference in the yield is largely attributable to the elimination of bank syndication and bank lending fees.

One last thing about what is private credit relative to other forms of credit, public credit, is that the credit rating, the credit opinion and assessment by credit rating agencies regarding that loan, regarding that security, is private. Meaning that it is available. The private credit rating itself

is available to the contracting parties and to the National Association of Insurance Commissioners (NAIC) and the rating rationale report is available to the contracting parties and to the NAIC. But it's not public on the website of rating agencies often and maybe most commonly at the preference of the borrower. So you might have smaller companies that are not public companies that are issuing debt and they need that debt rated and insurance companies want to buy it, but it is not beneficial to that company to have that information out available to the public and might, for example, include trade secret information. However, as my colleague can speak to better than I, having worked at a rating agency, under federal law, the methodology and every other aspect other than the publication of that rating is identical to a public rating. Rating agencies must use the same methodology for public and private securities, and they must put those methodologies out for public comment when they're developed and upon any change. So, I'll wrap up on what is private credit by just saying there is a lot of conversation about ratings and private versus public ratings. They're identical in the development and the output. The only difference is whether it's publicly available on a website.

Rep. Stephen Meskers (CT) stated private credit lacks clarity and valuation. It's not very well defined on a maturity schedule sometimes because private credit many times is financing to a bridge to nowhere, assuming there's a refinancing opportunity. Our insurance industry is probably one of the last bastions of security. The agencies did a horrible job during the mortgage crisis in rating publicly traded securities. Now you're talking about insurance companies buying privately rated, non-disclosed ratings to acquire debt that doesn't have a secondary market valuation. I can't understand why I would want my insurance companies to own anything that doesn't have an ongoing market valuation on a bid and offer and an idea of the spread and the risk. Because you are ultimately taking me to a bridge to nowhere and I think they can do that in their equity portfolio. They can take the credit risk that they determine in publicly traded equities or even some privately traded equities, but when I go into the debt markets and I start putting people into securities that if there is a rush to the exit, there is no bid and there is only offers. I can't imagine why I would want my insurance companies locked up in that business unless it's for a speculative portfolio. I just don't understand.

Caitlin Colvin, Managing Director of Insurance Regulatory Solutions stated I think it's a completely fair question. Maybe we can bucket the two points. It sounds like one is valuation, and one is more what are we financing and where is the off ramp? Rep. Meskers stated and liquidity. Ms. Colvin stated I think we can speak to what Blackstone does and a lot of times what we're what we're talking about when we're talking about what Blackstone does is what we view as a best practice. So, let's start maybe with that sort of bridge to nowhere point and what we typically are financing at Blackstone on the private credit side of the world. I think what you are talking about is that sort of private direct lending bucket - that lending to middle market sponsor-backed companies. Blackstone has 20 plus year's experience investing in that type of product and I think when you hear our COO and President talk about this, you talk about the underwriting and how important it is. We are not talking about credit ratings at all in this case. We are really talking about the underwrite and where are you on the loan to value? And this really gets to your valuation point. And Ms. Hagan and I sit on all of our investment committees and we think about is this an asset that makes sense for insurance companies to hold. And in many cases, those direct loans are not appropriate assets for insurance companies. You're absolutely right. They're typically rated non-investment grade in and of themselves. However, they're also typically very low value particularly when you look at a track record of a company in the higher and upper middle market range that really has a track record. I think there's also a distinction across companies of who you would lend to as well. So, when you're thinking about what we do, you're looking at senior secured loans at about a 30% loan to value and we can talk about that value and a lot of this has come up recently in the software space, for example.

Those software companies were overvalued, and a lot of that led to a handful of news stories about some idiosyncratic issues. But that's just what we're seeing in that market, it's pretty idiosyncratic. And the minority of what we do is that kind of thing and where we are lending to any software company, the loan-to-value is very low, and we are a super senior lender.

Rep. Meskers asked if the origination is to sell out 100% or do you retain a portion for your own risk portfolio? Because I think that's very important too. Because if you're originating and selling out 100%, the only person holding the bag is the owner and there's always an argument that should you have skin in the game. Ms. Colvin stated when we're talking about the direct lending, they're just debt investments. If there is a Blackstone equity component alongside of it, it's likely small for the alignment of interest point. But when we think about what we do, we are an asset manager so we are originating that loan for our clients to hold and take that yield.

Ms. Hagan stated that we are a fiduciary for our insurance company clients by law and by contract. The distinction between what Blackstone does and what banks do is that banks will originate assets and move them off of their balance sheet for a lot of reasons. We originate an asset, and the same team manages and monitors that asset through the life of the asset in partnership with our insurance clients. And that is actually skin in the game and alignment of interest that we think is really important. We are on the hook. We are only in this business as long as we manage assets carefully for insurance clients. We're not originating assets and moving them off of our balance sheets. We're responsible for their performance.

Ms. Colvin stated that's absolutely right and I do want to go back to the insurance company holdings point because I think this was really your point in terms of are these really assets that make sense for insurance companies? This is not the majority of what we do for insurance companies in any way, shape or form and in many cases, our insurance clients are holding this type of credit exposure to middle market companies through a securitization form. So, there's diversification in that pool of assets that are in that transaction and the insurance companies are typically holding pieces of the debt that are very highly rated and are nowhere near the first loss. So you're sort of holding that call it triple A to single A rated and I know your point on the public ratings. I get it. Rep. Meskers stated I think we got those triple A ratings on the mortgage portfolios. Ms. Colvin stated you do but the NAIC now rates all of those and totally, you can't rely too much on the ratings, but you can rely on the underwriting. And that's really where we're telling our clients and how we're explaining this stuff to our clients. We're not saying "hey, here's your triple A rated thing, take it and go home." We're not doing that at all. But I will say most insurance companies are not holding these things individually directly on their balance sheet. They're holding them in a pooled form, in an entrenched form and not the first loss piece.

Ms. Hagan stated that stated one quick thing I wanted to add is this piece of the conversation I think sometimes gets lost - we talk to our clients all the time about asset liability management and about the reason that these assets are important to reducing costs for consumers and increasing the availability and the affordability of insurance. It is very important to outcomes for consumers that insurance companies can invest in high-quality assets. On the illiquidity point, which gets back to the liability side of the balance sheet, we help our clients run cash flow testing. Every insurance company is doing that obviously; you need to do that. You can't be using your illiquid assets to pay liquid claims. That's why, for example, in the health insurance and property and casualty space, private credit will never be as large a part of their balance sheet and we are very mindful of that. But it is appropriate in our view for companies with long-duration liabilities to be holding long-duration assets. That's good governance as long as everyone understands that they have more than enough liquid assets to pay claims. Which is part of how we support them in that modeling.

Ms. Colvin stated I want to come back to your evaluation question in a second, But I think when you think we have a slide on this in the deck that talks about what we typically do for our insurance clients, and where are we spending most of our time. These days, we are spending most of our time on private transactions that are infrastructure or project finance related. You're financing hard assets so you're looking at something like power generation for a data center complex for Meta or Google or one of these hyperscalers that's currently investment grade rated. And in financing those assets, particularly for insurance companies, what we are most sensitive to is that debt amortizes or pays down in connection with a contract that is required to pay no matter what happens like a triple net hell or high water lease or contract. And that is really where we're focused these days and the majority of what we're originating for our insurance clients and we think these are really appropriate assets for insurance companies that are helping finance the real economy and they're long-dated and contracted.

Rep. Meskers stated in the banking crisis and going into the banking crisis, the bank originators no longer had that fiduciary responsibility in origination and sale. Are you telling me the asset managers do have a fiduciary responsibility in origination and sale? Ms. Hagan stated I can't speak to the legal framework around banks as I'm not a banking regulator. Rep. Meskers stated during the crisis it was fairly common that they didn't have that. It was more of a buyer beware. Ms. Hagan stated that I can speak to Blackstone's situation, which is we are a registered investment advisor subject to fiduciary responsibility rules and by contract, we accept fiduciary status and eagerly. It's not a reluctant thing we do at all with each of our insurance clients. We take that fiduciary responsibility very seriously.

Ms. Colvin stated so we talked about two big buckets of private transactions. We talked a little bit about the direct lending and that's not a huge part of what we're doing for insurance companies so it wasn't a major part of the slide presentation but really appreciate you asking because it's a really important point. So we talk about this bucket of infrastructure assets and financing the real economy. As you all know, and from your states you see it, there's a massive capital expenditure (capex) boom across all things relating to AI. We just had a really great AI presentation and that needs power, that needs data, and in order to do that, that requires a big build and a whole bunch of capex. So, those transactions that we're looking to help insurance companies invest in and be able to help finance that real economy at a little bit of a higher yield from public bonds, but at the same or better credit quality plus contracted cash flows and security over a real asset. We view this as a really attractive asset class for insurance companies and we're seeing a lot of these deals come in just given where everything lives today and what the world needs financed. Insurance companies are just a part of that. And then the other big bucket that I would say of assets that we look at on a private basis are private asset backed securities and we already talked a little bit about this, private middle market collateralized loan obligations (CLOs), which are those direct lending loans that are in a pool.

So, we have a little bit of information around what is the difference between a public and a private securitization or an asset backed security transaction. The answer is nothing other than the fact that the rating itself is more than likely private and not publicly available on Bloomberg or publicly available on the rating agency website. But both public and private securitizations have the exact same structures. They have the same type of collateral in them. Now, what collateral you focus on is very important when you're thinking about an insurance company. At Blackstone, when we're talking about securitizations, what we're typically talking about are high FICO secured consumer loans or residential mortgage loans or anything else financing hard assets or the real economy. We love doing that. So, any of those types of assets, typically in loan form, you can take those loans and pool them together and receive a diversification benefit.

If one loan goes wrong, you still have all these other really great loans that are still paying and you have a capital structure in the debt that is protective and most of the time insurance companies sit at the top of that capital structure. And this is just an example just from creating the loan and originating the loan to putting them into a special purpose vehicle (SPV) to essentially issuing different tranches of debt, and those insurance companies stay at the top of that debt capital structure.

Ms. Hagan stated I want to add one point that isn't related to any of the specific asset classes that we discussed, which is with private credit it's important for everybody to understand this asset class why it's on insurance company balance sheet? Look under the hood. It's important for insurance companies to know what they're investing in and to be really comfortable in terms of underwriting, etc. I will just say that private credit is really not new and you've probably heard this, so forgive me, but a great example of it is that a lot of private real estate transactions in this country were insurance company lending to build things like the Empire State Building which was financed by MetLife. And one other last point I would make on a related note is that we are not obviously exclusively engaged in private credit lending with insurance companies. In addition to other asset classes, we have a lot of strong partnerships with sovereign wealth funds and public pensions specifically in the states, and that's something that's gone on for many years. Again, we do a lot of work with the states, including across asset classes but also including private credit. It's just those types of investors haven't been so much at the forefront of the conversation but I wanted you to be aware that they're very much part of the ecosystem.

Asm. Gandolfo stated this was a great presentation and this was 101 style and it does sound like there is interest in this and we might have to revisit this topic again in the future.

DISCUSSION ON FEDERAL EXECUTIVE ORDER "ENSURING A NATIONAL POLICY FRAMEWORK FOR ARTIFICIAL INTELLIGENCE"

Asm. Gandolfo stated now we will turn to our final presentation, which is a discussion on the federal Executive Order (EO) titled, Ensuring a National Policy Framework for AI. As a reminder, NCOIL did recently adopt a resolution affirming the U.S. state-based regulation of AI in insurance, which is consistent with the McCarran-Ferguson Act. Since the passage of the McCarran-Ferguson Act, states have consistently exercised primary authority over the business of insurance in a manner that protects policyholders and claimants, ensures insurer solvency and promotes fair, competitive, and stable insurance markets. Against that backdrop, NCOIL and nearly everyone interested in the state-based system of insurance was troubled by the issuance of this federal EO that essentially aims to preempt states' ability to legislate and regulate AI. Whether the EO can actually preempt states is an interesting question and we'll hear today information about that and about the EO and the practical implications for the states.

Oliver Engebretson-Schooley, Partner at Sullivan & Cromwell, thanked the committee for the opportunity to speak and stated that I will preface that with I am not an AI expert, so if you have technical questions about AI, we're going to have to bring Mr. Armstrong back in here but I can talk to you about what this EO means, and probably more importantly, what it doesn't mean. I won't bury the lead here. I'll get into this in more detail, but does this EO preempt state law? No, it does not. And that's true for any state law. Not just the state laws regulating insurance although as Asm. Gandolfo mentioned under McCarran-Ferguson, as many of you likely know, state insurance regulations are specially protected from federal preemption. So then, what does this EO do? It's really a roadmap. It lays out the administration's policy on the regulation of AI, and then it provides a roadmap of strategies for how the federal government might go about implementing that policy and making it a reality. And it's some of those strategies

that I think are actually much more interesting, and that I'll talk a little bit about today. Some of them are already being implemented beyond the preemption question and we'll talk a little bit about what that means for state regulation of insurance. The EO was signed in December 11, 2025 and naturally spurred a lot of hand wringing from folks who from state legislators and regulators who are interested in regulating AI, it says it right there in the beginning the purpose of this EO right at the top is to create a national framework for the regulation of AI and limit states' ability to regulate the industry. Again, as I mentioned at the top, this EO does not preempt state law. It is purely precatory. It does not purport to preempt state law, and in fact it cannot, and I'll get into why in a moment. It doesn't actually impose any binding requirements at all on states, on developers, or on users.

Just taking a step back briefly, what is federal preemption? How does it work? So, this is a doctrine that's founded and grounded in the U.S. Constitution's Supremacy Clause which states that the laws of the U.S. shall be the supreme law of the land and courts have interpreted that to mean that Congress can enact laws that preempt state laws. So, Congress can do this in a couple of ways. It can be explicit about it. Like the federal Employee Retirement Income Security Act of 1974 (ERISA) explicitly preempting state laws that conflict with ERISA is a good example of that. It can also do it impliedly where courts interpreting statutes decide that Congress's intent, even though they didn't say it explicitly, was to preempt state law. And this can show up in what's called field preemption. So, that's when Congress passes legislation and then there's additional regulation so pervasive in a space that courts just say, look, Congress clearly here was intending to occupy this entire space, occupy the field, and states just can't go there. You have no authority to regulate in this space, regardless of whether there's any actual conflict between a state law and a federal law. States just can't go there. The other form is conflict preemption. This is much more common. That's when there is a conflict between a state law and a federal law. So, requiring someone to do one thing when the federal law requires another, or even when a state law is determined to prevent the objectives of the federal law, they can find preemption there as well.

But the key point to take away from all of this is that preemption requires Congressional action. But, to be clear, that does not mean that Federal regulations cannot preempt state law. Federal regulations are law, and agencies are creatures of statute. But they have to be acting within congressional authority that has been granted to them. They cannot just go out of their way and say, "I want to now regulate this space." And the Supreme Court has actually said that even sometimes EOs can preempt state law but in those narrow contexts, it's when Congress has said, "White House, you can regulate this particular space." And so, the White House has congressional authority to issue regulations that then preempt any state regulations in the same area. An EO cannot on its own without any sort of congressional authority, nor can an agency, just say we are now preempting this state law. So that's the general framework for preemption. As I mentioned, and again, as many of you are aware, state regulation of insurance is unique because under McCarran-Ferguson, states are actually given the benefits of what is sometimes called reverse preemption or anti-preemption. So, when a state law is specifically enacted to regulate the business of insurance and the relevant federal law that has some sort of conflict is not specifically to regulate the business of insurance, the state law controls.

So, let's turn back to the EO and its key sections, which I've summarized up here on the slide. So, section two sets out the federal policy. As I mentioned, it's to sustain and enhance the U.S. global AI dominance through a minimally burdensome national policy framework that would displace what the federal administration has called a patchwork of state laws. Section three then directs the Attorney General to set up a litigation task force to challenge in court state AI laws that it views to be inconsistent with the policy laid out in this order. This one's particularly

interesting, and we'll talk about it a little more in a moment. Section four and five, this directs Commerce to evaluate current state AI laws to identify ones they consider to be onerous and that conflict with the policy in the EO and then commerce can either refer those to the litigation task force or the EO also directs Commerce to come up with a framework for restricting certain grants and funding to states that have these laws in place. And then sections six and seven direct the Federal Communications Commission (FCC) and the Federal Trade Commission (FTC), two agencies that do have regulatory authority that has been found to preempt state regulations, and directs them to investigate "can you use these existing authorities to try to preempt state regulation of AI?" And then section eight says, "Let's come up with a legislative proposal that we can then present to Congress to actually put in place a federal uniform regulation of AI." So, then back to the key question that we started with what does this EO do? As I said, the key takeaway today should be that it does not and cannot preempt state law. It is just precatory and as I said, it does not purport to. It does set out executive policy on AI regulation and then provides this roadmap of strategies that I mentioned. At the moment, it does not seem that there is much appetite in Congress to create this national framework that was mentioned in the EO which would give agencies the authority to start issuing regulations that might preempt state law. And the White House can't do that on its own.

And again, even if there were congressional appetite at the moment to enact this type of broad legislation, under McCarran-Ferguson that still would not preempt states laws regulating the business of insurance unless that new statute specifically said this is to regulate the business of insurance. So there's that additional protection there even if Congress were to act here. And then I mentioned the FTC and the FCC. They do have authority to issue regulations that can preempt state law. They could, in theory, develop these regulations. So far, the FCC, I am not aware of them taking any particular action. A couple of weeks ago, the FTC did issue what it called a proposed policy statement where it was trying to lay out a framework for the FTC to try to preempt state laws regulating AI usage that would conflict with its section five authority to regulate unfair and deceptive practices. But again, these are just sort of patchwork attempts. And going back to McCarran-Ferguson, there's the additional protection from a preemption risk for insurance specific legislation.

But that does not mean that this EO should be ignored. There are a couple other strategies in here that are laid out that I think should be a particular interest to state legislators. The first is this evaluation of state laws and the directive to Commerce to come up with a list of laws that it views as contrary to the policy laid out in the EO. I am not aware of this list being public yet, that doesn't mean it doesn't exist in some form. What is clear is that federal funding is going to be a key tool of the administration to try to influence state policy. The BEAD (Broadband Equity, Access and Deployment) program was specifically mentioned in the EO but the Order also directs other agencies to examine discretionary grant programs to see if they can also be used as tools to try to influence state policy and regulation on AI. The far more interesting one, I think, and the one that I will spend the rest of this presentation on is the new litigation task force. It is up and running and it may well prove to be one of the most active tools of the administration, at least in the short term, to try to carry out the policies that are laid out in the EO. There are three main litigation strategies that they seem to have developed. Two of them are mentioned specifically in the EO. One of them has been pursued in a particular case. These are strategies for the federal government to challenge in court state laws attempting to regulate AI development and use. So, these strategies are preemption, commerce clause challenges, and equal protection challenges. We already discussed preemption. There is very little risk, if any, of a preemption challenge to a state law or regulation in the area of insurance, and that's actually true as well for the second strategy or dormant commerce clause. Just a quick primer on that. So, the dormant commerce clause under that doctrine, a state law can be deemed

unconstitutional if it unlawfully discriminates against other states, specifically if it doesn't have some sort of local purpose and if it imposes burdens on interstate commerce that are deemed excessive relative to the local purpose.

But again, under McCarran-Ferguson, because Congress has granted states the authority to regulate the business of insurance, Courts have said Congress has also granted states the authority to regulate the interstate flow of commerce. So not only is a preemption challenge off the table under McCarran-Ferguson for state regulations of insurance, but Commerce Clause challenges are as well. So, that leaves the Equal Protection Clause argument. Out of the three litigation strategies, if state insurance laws, at least at the moment, are to face litigation risks, from what we have seen, it's going to be from this one. So, the Equal Protection Doctrine stems from the Fourteenth Amendment, at least when it's applied to states, which prohibits states from passing laws that deny to any person within its jurisdiction, the equal protection of the laws. That means that state laws discriminating against certain classes of people are constitutionally suspect and as many of you are likely familiar, laws that discriminate based on what are called suspect classifications, like race, are going to be subject to strict scrutiny and will often be invalidated. As use of AI tools proliferate, many have been focused on ensuring that AI tools are not used to or do not perpetuate discrimination, one example of that is Colorado's since revised algorithmic discrimination law, which was designed to require developers to ensure that AI tools did not result in discriminatory impacts among a number of other things. These laws have become a target of the federal government and they face litigation risk, not just from the government, but also from private litigants, and I'll talk a little bit about a recent example of that. So, this actually involved the Colorado law that I just mentioned. This statute was passed, it sought to regulate AI use in what were deemed high-risk areas. High risk areas like employment, housing, credit, health care, and insurance. The law required AI developers and users to exercise reasonable care to prevent discrimination.

And specifically in the law, that included what it called disparate impact on protected groups. So, what does that mean in the discrimination context? Disparate impact is talking about even if the law or a policy or a tool does not intentionally or expressly discriminate against a particular group, controlling for other factors, the policy or tool ends up treating that group differently. It's not a subjective test. It's an objective one that looks at the outcomes and if there are disparate impacts, the Colorado law required developers and users to take certain actions to correct for that. xAI sued in federal court to block enforcement of this statute, and then shortly thereafter, the federal government intervened through this new litigation task force that was set up through the EO and it brought a Fourteenth Amendment claim saying that this law violates equal protection. The federal government argued that by requiring developers and users to prevent the risk of disparate outcomes based on demographics like race, it effectively required them to intentionally use those demographics when building and running the models. And that's a problem, the federal government said, because particularly in what are called zero-sum games, so something like a scenario like hiring where only one person can get a role, if you are taking some action to try to improve outcomes for one group, you are necessarily discriminating against another. This was a concept that was discussed at length in the Supreme Court's affirmative action decision, the Harvard admissions case, that came out a few years ago. If you try to do something to favor one group, you are going to disadvantage another and that's unlawful discrimination. The Colorado case was dropped fairly shortly after the federal government intervened because Colorado rescinded the law. It has since passed a new one that does not include much of that same language that was being challenged in this case. So we don't know how this theory will fare in courts. I have not seen it come up since, but I have little doubt that we will see it soon as states continue to legislate in this space and as the federal government's litigation task force gets further up and running.

So to bring this back to state insurance regulation, what does this mean? For one, insurance regulations like the Colorado law that impose duties on developers and users to try to prevent discriminatory impact could face similar legal challenges from private litigants and from the federal government. Many states already have these types of regulations in place. Now, insurance is not hiring, it's not necessarily a zero sum game. And so there are some basis to argue that it should be treated distinct but in the affirmative action context, where this was being discussed in that case, there already was a recognized interest in universities to create a diverse student body. The Supreme Court has already said that just remedying societal discrimination generally is not a compelling interest and so it's likely that if states continue to try to issue regulations that require developers to make these types of changes that they will see additional challenges and courts could be receptive to them. As I said, though, the case law here is very much still developing. I suspect that we will see litigants, both private litigants and the federal government, pressing this broader theory that any state law that requires developers to favor a particular group, regardless of whether it's a zero sum context, is unconstitutional. But at a minimum, I think state legislators should continue to monitor the legal landscape here so that they know and understand the law as it currently stands and what the risks are as they continue to develop state regulations for AI and the use of insurance.

Asm. Gandolfo stated you mentioned before about a scenario where Congress does decide to step in and pass some kind of law to establish the federal government as the sole regulator of the use of AI and the development of AI. Now, with McCarran-Ferguson in place, would this hypothetical bill need to specifically reference the use of AI in insurance? Or would a broad encompassing attempt at preemption cover it? Mr. Engbretson-Schooley stated, no, my understanding based on McCarran-Ferguson is that if they were to enact a new statute establishing this uniform framework, if they wanted it to specifically preempt state laws regulating the business of insurance, even if in the AI context, they would have to specifically say that is what this law is doing. They could do that, but that's an additional step that they would have to take.

ADJOURNMENT

Hearing no further business, upon a motion made by Sen. Lang and seconded by Rep. Lampton, the Committee adjourned at 5:30 p.m.