

CTL-CISO-Billy-Spears

📅 Mon, Jul 13, 2026 6:34AM 🕒 16:22

SUMMARY KEYWORDS

AI governance, operationalizing AI, legal accountability, identity as control plane, policy enforcement, shadow AI, third-party MCP servers, zero trust access, behavioral monitoring, supply chain provenance, strict delegation models, action-level authorization, runtime authorization, AI risk.

SPEAKERS

Billy Spears, Jo Peterson



Jo Peterson 00:07

Hey y'all! Thank you so much for joining Clear Tech Loop. I'm Jo Peterson. I'm the CIO of Clarify360 and the chief analyst at ClearTech Research. And I've got a treat for you today. I've got mr. Billy Spears. Hi, Billy.



Billy Spears 00:22

Hey, Jo. Hey, everyone.



Jo Peterson 00:24

Hey, thank you so much for joining. Billy is the CEO of a stealth startup in cybersecurity, so more to come on that. Billy is a cyber and IT veteran with time spent in executive roles at companies like Dell, Hyundai, LoanDepot, where we met way back in 2018. He has served as an adjunct professor of cybersecurity, and I'd like to take a moment to thank Billy for his service as a U.S. Marine.



Billy Spears 00:55

Thank you.

J

Jo Peterson 00:57

You're welcome. So, as you guys know, we do three hot take questions, and Billy's in the hot seat today, going to answer the questions. But the first one got Billy. How do we operationalize AI governance, and who is legally accountable when an AI agent makes an unauthorized decision? What do you think?

B

Billy Spears 01:17

Yeah, I think I think we'll start. We'll make it easy, right? AI governance is not a policy problem, but the uncomfortable truth is most companies are not operationalizing AI governance. They're writing policies and calling it governance. Real governance shows up in execution. Operationalizing AI governance requires three things. First thing is control placement across the AI lifecycle. For example, you have data ingestion. So, folks at home, that's provenance, integrity, validation. You have model training. So, what do I mean by that? Isolation, tamper protection. You have inference, runtime controls, policy enforcement. You have agents. That's action level authorization, not just prompt level filtering. This is exactly the the models that I've implemented across enterprise AI platforms. Number two, you have identity as the control plane. Every AI action must resolve to an identity, not the model did it because that's what we're hearing today, right? We're like, hey, I got a model. It does stuff. Woohoo! We're good. That's not true. It's the system acting on behalf of this user under this policy. If the identity is not embedded, governance is just fiction. But we're hearing it, right? AI does the thing, so we have governance because we're following some checklist, and lastly, policy enforcement. It's not policy documentation. Governance must be machine enforced. Think runtime authorization engines, not PDFs. And then, lastly, I think this is really important to the spirit of the question, Joe. Who's legally accountable? This is where the the market is still pretending, and I know that I'm being right on the edge. I'm not really trying to go left or right here, folks. But you need to hear it. The AI is not accountable. The vendor is rarely accountable beyond the terms of service. The enterprise is deploying the AI. The AI is the one that's actually accountable. Specifically, let's be like really laser focused here. The board and the executive leadership team own the risk acceptance model. So if you're not talking to them, that's the first place. Go talk to them. Your C team, CIO, CSO, CTO-they own the control design and the enforcement, but that's higher level. So that's like architecture, right? Product owners own how AI is mentioned. So if an AI agent makes an unauthorized decision, regulators aren't going to ask, "What did the model intend? They're going to ask what controls failed and who approved the system. The bottom line: AI doesn't change the accountability; it accelerates how fast you can get it wrong. So think about the holistic life cycle of the question, and I and I hope that I got the spirit there because really, when that model makes the decision. It's not really, hey, do I have the documentation? It's do I have the life cycle of governance solving the equation?

J Jo Peterson 04:31

You got it 100% right, and I wrote it so that it was a little spicy because right on purpose because this is going to fall right back into Soso's lap, goes right back in his or her lap because it's part of the risk equation. At least in my mind, it is. So, just it just is. I recently heard MC. He referred to as a confused deputy. I just I love that. I don't know why. How do we prevent agents from executing actions that the user should not be allowed to perform?

B Billy Spears 05:15

Yeah, I think the confused deputy framing is accurate, and it's it's understated, Joe. You know, most AI agents today are overprivileged middleware at best, right? Their inherited authority without enforcing intent, and the the the failure pattern kind of. I'm going to walk your users kind of through another series of things, right? It kind of looks like this: the user has limited access, the agent has broad API access, right? Agent executes an action on behalf of the user. There's no enforcement of user level authorization at the execution time, right? That's not an AI risk. That's a basic security failure. So what's the check for for security? So hey, security teams, what's the check that you have for for your development teams to set up these MCP servers, right? So then, when you get to the fix, right? What's your strict delegation models? Here's what I think. Free. [Here's some free advice from Billy and Jo, right? Strict delegation models. Agents must operate with constrained delegation tokens. Must not system level credentials.](#)

J Jo Peterson 06:27

[Yeah,](#)

B

Billy Spears 06:27

no implicit privileged inheritance. All right, agent or excuse me, action level authorization checks. Every action must be revalidated at execution time, not at the the prompt time. This is where the the implementations fail. Policy binding to intent and identity. For example, who asked, plus what is allowed, plus what action is being taken. All three must align before you allow execution to occur. That's the governance. of your policy. So your policy is actually encapsulated or on the wrapper of some actionability. And then, lastly, kill the helpful agent anti-pattern. Meaning, if the agent can figure it out beyond the policy boundaries, you've already lost. Meaning, if the technology can outthink your governance strategy, it's terrible, right? And I know that's a little controversial, but it's accurate. Most organizations are deploying agents faster than you've deployed role-based access controls 10 years ago. We are repeating the identity mistakes at machine speed. And 10 years ago, we asked you this simple question: How many actors are on your network? And every one of you, as we went around the country, gave us all kinds of answers: 210, 200 507,000,003 But the answer is two. There's humans and machines. We spend \$12 billion on humans and give you cards and access and authorization tokens and whatever. But we let as many machines as the world in the world get on inside your network. And I'm telling you, Skynet is here, people. So every AI action has to resolve to an identity. Not the model did it. The answer won't survive a regulator or a breach review,

J

Jo Peterson 08:26

right? And we haven't even, and that was such a good answer. And thank you for that. You know, my head starts spinning when I think about ephemeral agents and the permissions that are assigned to ephemeral agents, and who's who's governing those, right? I mean, it just it it's a little mind boggling. Yeah. Third question: How do we verify the authenticity and security of third-party MCP servers?

B

Billy Spears 09:03

You guys may not like this one, but my short answer is, don't trust them.

J

Jo Peterson 09:08

Right, you have to

B

Billy Spears 09:09

continually verify them. Right, I think the the longer answer is right now most organizations treat MCP servers like trusted extensions.

J Jo Peterson 09:18
Exactly,

B Billy Spears 09:19
they're not. They are untrusted execution environments with privileged access paths.

J Jo Peterson 09:25
Yep.

B Billy Spears 09:27
So what needs to really happen? Okay, folks, sit down, sit in your chair, buckle up your seatbelts, and let's get into it. You got to have a strong identity and attestation here. Cryptographic identity for every MCP server, you got to have signed workloads, hardware-backed attestation where possible. Number two, zero trust access model. How long have security people been saying zero trust? Not some trust, maybe trust. We think we can trust zero trust, no implicit trust based on network or vendor. Full stop. Every request has to be evaluated in real time, and you have to have least in privilege enforced at the API level. Full stop. Number three, you have to isolate by design. Then we're talking about sandbox execution, no lateral movement, no direct access to sensitive systems without mediation. Whoa! I know. I just got crazy, crazy in here. But this is serious business. We're talking about stuff that moves faster than you can think. So we have to get back into number four, which is behavioral monitor. Excuse me, behavioral monitoring. You are not just validating code. You are validating behavior over time, meaning more information than your body can process, more information than your team can process, more information than your company can understand in any simple amount of time. This is AI-driven anomaly detection, so it is critical that you are monitoring this behavior because it is not sufficient alone without the validation.

J Jo Peterson 11:06
Yeah.

B

Billy Spears 11:06

So you have to extend it out to your supply chain. You also have to validate your supply chain provenance of models. Yeah, we're talking about models, plugins, dependency, and for those of you who don't know what I'm talking about, your S bomb equivalent for AI services. You have got to get that done because listen, most third-party AI integrations today would not pass a basic third-party risk assessment if evaluated honestly, but they get a pass because they're labeled AI. And if you don't believe me, you talk to your DevSecOps team. You talk to them right now, and if you don't have one of those teams, you talk to your AppSec team and you ask them to go do a check on them right now. And if you don't have those people, you go talk to security engineers. I guarantee you, you're gonna find three things: AI governance is not a policy problem; it's execution. AI risk is not a new risk. It's poorly controlled existence risk at scale. And agents are not intelligent. They're just automated decision amplifiers. Companies that win AI will not be the ones with the most models. They'll be the ones that can prove control, trust, and accountability at runtime. We've been saying this for years. You guys gotta do it. We're talking about cloud security at its finest here.

J

Jo Peterson 12:26

Yeah, I mean you're right. And what lessons? Such a good, valid point because I feel like sometimes it's Groundhog Day. I remember when I really got into architecting cloud workloads, and I remember going, "Where's the security? And they're like, "Oh, we're gonna put it on after. I'm like, "What do you mean you're gonna put it on after? Well, yeah, we put that on after. I'm like, "That doesn't sound right. That's not good. It's not good at all. But I feel like we're in that same place with AI, and I have these conversations about with folks, and I'm like, "What are you doing around shadow AI? Yeah, that's a problem, right? That's that's what I get back. Well, we're thinking about it, okay, right? So it's it's more of an issue than I think people think, and/or you're not sure how to approach, or but it feels like the early days of cloud again to me, at least.

B

Billy Spears 13:31

Yeah. Well, I think this is interesting. So I like how we're we're kind of pivoting a little bit. Shadow AI is the modern version of of shadow IT. It's just far more impactful and a much more it's a it's a much shorter blast radius, right? And let's be honest, it's already everywhere, right? Because companies are pivoting. It's a new it word. It's like a verb or adjective that extends top line revenue, and it's already everywhere, right? So employees are using AI agents are using plugins, copilots, external tools to move faster, smarter, whatever. But really, what are they producing outside of the AI companies? What are companies actually producing that's driving value? Here's what I think they're doing: they're summarizing, they're summarizing customer data, they're generating some code at what 1020, 30% rate, and then they're validating. They're drafting contracts, analyzing data sets, often outside of any approved environments. At this point, if a company's out there that's hearing this podcast that wants to refute it, hey, contact Joe and the the Clear Tech Loop, and let's get back on this thing and let's debate because I really want to go at this thing because I disagree with you. Like, let's show me that I'm wrong because I think you don't stop with policy. You don't even slow it down that much. So the question isn't how do we prevent shadow AI. I think the question is how do we make it safe enough while still enabling the business.

J

Jo Peterson 14:56

That's it. That's see. That's the jelly right there. You nailed it, right? I mean, and and we're getting there, and I think we're getting closer. And as always, thank you for your thinking. I love that you spice it up a little bit and ask provocative and answer in provocative ways. So that's lovely, and I have to have you back again. So I hope you'll come and visit again.

B

Billy Spears 15:22

So I'll leave I'll leave you and your audience with this first. Let me thank you, Joe, and thank your audience for for listening to a few minutes of of time they can to get back in their life. I think in this broadcast, AI is not eliminating risk; it's amplifying the consequence of weak controls. We'll stop. The companies that win are not the ones with the most AI. They're the ones that can prove control when it matters, and it does matter, folks. So if you want to know more, contact Joe. If you want to debate me, again, contact Joe. Let's get on here and and have a debate. I'm sure Joe can facilitate that. I would love to hear what you have to say. If not, drop a comment below. You want to get a hold of me? I'm on LinkedIn. I would love a like and follow because I love that stuff. And and let's get into it. Thank you so much, Joe. You're amazing.

J

Jo Peterson 16:09

Thanks. Thanks for coming. Bye, y'all.