

How WongDoody keeps humans in the lead with AI | Box AI-First Podcast EP 20

23 MIN · YOUTUBE · [HTTPS://YOUTU.BE/QTFC0FCWz5k?SI=TVJLLZLVV0APTVID](https://youtu.be/QTFC0FCWz5k?si=TVJLLZLVV0APTVID)

<https://youtu.be/qTfCofCWz5k?si=tvjLLZLVVoaPTvid>

SUMMARY

The podcast discusses the challenges and considerations of implementing AI in enterprises, particularly focusing on security, governance, and responsible usage. Jeff Chambers from Long Duty emphasizes the importance of a robust strategy for AI integration, highlighting the need for security measures, privacy considerations, and the evolution of governance as AI technologies advance.

- Security is the first major concern when integrating AI into enterprises, followed by privacy and governance.
- Organizations often underestimate the time and return on investment from AI initiatives.
- A zero trust environment is essential for AI deployment to mitigate risks.
- Governance processes must evolve to keep pace with rapidly changing AI models and technologies.
- AI credits should be treated as R&D expenses, with careful tracking to prevent uncontrolled spending.
- Companies should focus on restructuring processes around user permissions and data access to protect sensitive content.
- The concept of "human in the lead" emphasizes the need for continuous human oversight in AI operations.
- Ethical considerations in AI deployment, including environmental impacts and local economic effects, are becoming increasingly important.

Speaker 1: The very first thing that typically breaks is security. And that's just because it's such an emerging technology. And while it's been here for a while, everyone has been able to use it on a personal capacity. Right. People say ChatGPT is the very first thing people think about, and the very first experience they have is either a mobile app, right. Or it's a web browser with a prompt to that regard. People are not thinking about security on the personal level, and then they try to use this for experimentation inside the enterprise.

Speaker 1: And that is where the real challenge is.

Speaker 2: This is the AI first podcast hosted by me, John Herstein, Chief customer officer at Box. Join me for real conversations with CIOs and tech leaders about reimagining work with the power of content and intelligence and putting AI at the core of enterprise transformation. Well, Jeff, welcome to the AI first podcast. Let's start with the quick introduction. Tell us a bit about Long Duty yourself and your role there.

Speaker 1: Well, my name is Jeff Chambers. I'm the VP of IT Technology based in Los Angeles, and I work for Long Duty, which is the global creative technology company and arm of Infosys.

Speaker 2: I want to start by just asking you, given the purview that you have when AI starts getting used across the organization, what are you seeing that breaks first? It's kind of a funny question, but is it security? Is it cost control? Is it accountability? Maybe it's maybe something not on that list. What do you see breaking in kind of the first instance?

Speaker 1: Well, the very first thing that typically breaks is security, and that's just because it's such an emerging technology. And while it's been here for a while, everyone has been able to use it on a personal capacity. Right. People say ChatGPT is the very first thing people think about, and the very first experience they have is either a mobile app, right. Or it's a web browser with a prompt. So to that regard, people are not thinking about security on the personal level.

Speaker 1: And then they try to use this as for experimentation inside the enterprise. And that is where the real challenge is. So security is the first thing that breaks. And then of course, after that, it is privacy and the foundational structure and vetting of how AIs use the organization.

Speaker 2: You know, the way you described, you know, your role in the organization, you're not directly leading the. The AI work for Creative, that, that, that sits somewhere else, but you are responsible, at least in part, for governing it. Right? So where does the tension show up between getting that work done, allowing people to do the things that they need to do and then putting the right sort of guardrails and governance around it.

Speaker 1: Right. Governance is I think the hottest topic right now in AI. What we're seeing right now is I think what twice a week now we're seeing press releases from anthropic about models going faster than what humans can actually patch on systems, exposing vulnerabilities. So the governance that we have in long duty follows our parent company Infosys, which is an ISO 42001 certified company. And we are trying to vet each model while also allowing for experimentation inside of a controlled environment.

Speaker 1: The governance right now is always evolving and there's not many very pervasive governance body like models out there or systems that will govern the AI. So a lot of it is we would ask employees to bring their interest of the AI models. It could be a raw model. How do you want to run it, what is the platform that's offering it? And then we will run it through security, privacy and due diligence for those vendors to use it in the system and if they are allowed to use it.

Speaker 1: We are never going to scale anything unless it goes through three or four reviews to make sure that we are trusting that model at that time. Because we live in a zero trust world and we have to abide by that. So it is very difficult. And that is kind of the tension with the governance is that ability to use AI, be agile, be that fast paced company without exposing our risk and having that strong security posture.

Speaker 2: If you think about how you're actually governing AI usage today, sort of curious, as specific as you can be, what's allowed, what's restricted and how do you enforce that?

Speaker 1: Great question. For each geographic unit is a little bit different. Right now we are not allowing AI in any of our on any of our content which is primarily either in box or in SharePoint or locally in Germany or Serbia. And the models that we are using are custom models downloaded run locally either in in Europe on that content. For those clients that are not able to use AI, even running in a private cloud or directly from the manufacturer, the the governance there goes through of course the same types of channels and keeping that AI directly on that content in siloed with limiting access to it.

Speaker 1: In the US we are using box. And so what we are doing is we are have expanded our AI adding box enterprise advanced so we can utilize the 10/enterprise AI tools in there to be released on our data in a very, very deliberate way. And we are vetting the models and we only Pick models that are running generally on private data centers that are hosted by Box. So that that is what we're starting with.

Speaker 1: And then we are at the point where we are now then bringing those to a responsible AI team in India and vetting through the process of making sure that these AI models are cleared.

Speaker 2: So how are you thinking about AI credits, usage, tracking and really making sure that these experiments don't turn into uncontrolled spend?

Speaker 1: AI, just like any other tool, has cost. What are the costs? And, and even if I was to look up the cost today, or someone viewing this podcast, a year from now, they may not know what the costs are for something and cost is really irrelevant. It's almost like the Chuck E. Cheese example. How many tickets does it take to get that stuffed animal? Well, it changes. So when it comes to credit usage, we want to look at giving people the ability to you to do whatever they want to do in a safe environment, generate whatever element it is, whether it's designed to code, whether it's cleaning up code, whether it's summarizing documents, responding to RFPs.

Speaker 1: And the AI credit should be an R and D expense that is not actually in it, but is more linked to a OP X expenditure. And that is, that is what we're trying to push for, is that experimentation. Find out what you get from the AI credits. And there are of course AI credits do various things depending on the model, with all the variation and models that change every three months from the biggest ones, from OpenAI to Cloud to Watson, and the list goes on and on and on.

Speaker 1: The credits will give you different responses. So really bringing that in to run the same type of query response, analysis or generation against different models and seeing what you get is really at the core of where we're at now and many companies are at. So for now, we're not worrying too much about the credits. We are monitoring them as platforms will give visibility to them. Whether it's one of our leading platforms, which is figma, Adobe, each vendor is now, you know, building a what essentially should be a cost dashboard usage, who's using it, what it's being used against, how many files it's being used against.

Speaker 1: And so all these KPIs and metrics will come out of those dashboards to give a better picture about how much money is being spent on AI, how much is experimentation versus how much is being put into production. So that is really the biggest thing that I don't see people talking very much about right now.

Speaker 2: What would you say most organizations are underestimating about AI right now they're underestimating,

Speaker 1: I think the return that they're initially going to get in the time it'll take to get that return. It's a buzzword, right? We need to have AI first, right? We are an AI first organization, not just Long UD and Infosys, but many other organizations. So what is lacking is the strategy for AI, how to get there and the milestones and the whole life cycle of experimentation, right. Bringing it into a staging environment, vetting that by privacy, security, ethical nature, cost, and then running a feedback loop that will then make that work for a sustainable future.

Speaker 1: It's not a project that is evaluated once and it's done. It's ongoing change. And the change management around AI is just, it is a whole. Another beast to deal with. Now we've seen this within the last six months is we've recommended that use the model that's available available to you at that time. Do not try and customize any model specifically trained on any, you know, client data or customer data, because the model that's going to come out in three months is going to actually be able to do what you're asking it to, to do in specific task.

Speaker 1: Now if there is of course is a very specific use case for text generation or you want to train a model to make something like the data that you're giving it, then yes, there, there are specific use cases for that, for, for brand, for let's say medical, for technology. But there has to be a very specific use case to actually train a model and use that model just for that purpose.

Speaker 2: So I want to ask you, for your peers out there, other CIOs, if you had to really break this down and simplify it for another CIO who's trying to put together a strategy, what's the first control they need to be putting in place before they start scaling? So when you move out of experimentation and you're getting ready to move something into more of a production mode, what should they be thinking about from a security and compliance control perspective?

Speaker 1: So the first thing is that strategy change champions in each department, those knowledge based workers are going to be the ones who are going to be building your agents. Then of course you need to get some financial backing and set some metrics or KPIs on what you're trying to get out of it and build those use cases. Then of course zero trust environment is what you should be running these things in, whether it's in a staging environment or production environment.

Speaker 1: Yeah, right. Because you do not want to introduce any issues with a new model or even an old model giving you security vulnerabilities or risk to any of your clients or your own reputation for running, for running as a, as a business that's supposed to have AI integrated into their internal processes or for their customers. So the AI strategy is the biggest thing and just keeping that strategy going and evolving partners outside of the technology and data sets.

Speaker 1: You don't need to be a data scientist to do all of this. You need the knowledge based workers who

are, who know the processes. Processes need to be reevaluated from the ground up and so understand your complete business process across all of your systems where your content is and get that house in order.

Speaker 2: How are you thinking about protecting content, specifically sensitive content in this sort of AI powered or AI enabled environment?

Speaker 1: So the most important thing is to redo your processes, where data lives, the permission structures, all the foundational basics that all companies should redo. The strongest companies that are deploying AI in the enterprise are companies that have gone back and restructured all of the processes around. User permissions, systems access, API calls. Those are the strongest. If you have a very strong foundation for access the most relevant content, you know, whether it's confidential contact or content that is internal or contains PII or other highly classified information, restructuring that rethink it.

Speaker 1: Is the foundation to be able to use an AI agent? Yeah, because as we're seeing now, whether we're using Microsoft Co Pilot or whether you have anthropic or open AI integrated into any content system, that agent will typically run with the permissions of the person who's running them if they have access to the content. The AI is going to do a great job, whether you like it or not, finding the content. Is it relevant, is it not relevant or is it now finding content you never would have thought that exposed data and it might even be handing it off to another agent.

Speaker 1: If it is based upon user, the user running that permission, just like a basic permission model, then that needs to be strengthened. So I do see guardrails around foundation guardrails on an overarching agent orchestration security privacy compliance model that's looking at all the calls coming in and out. Right. We have API gateways for each of our products. Almost every company does. How are we protecting those? We want to ensure that vectors of attack, right.

Speaker 1: Or vectors of internal processes are all known what is coming through those interconnected systems. And we can do the same thing with, with AI agents. Whether it's an API call or an MCP server. So it is all based around restructuring user permissions, redoing the models of content and then protecting all the gates to all of your access and keeping that in a continual monitor loop.

Speaker 2: Are you seeing that definition evolve? Are people thinking about this concept of zero trust differently now?

Speaker 1: I've only heard about it recently in what I've read and what I've seen. I've seen nothing in practice per se, except for what Microsoft maybe what other vendors that are providing enterprise AI where you will add a layer of kind of AI governance for looking out for some of the most basics which are PII or any of the regulatory things like exposing Social Security numbers or EINs or checking at a routing number. So those things I have seen in the box platform and in a few other platforms, but as far as each agent having guardrails inside of it, I have not seen that.

Speaker 1: That is something I would like to see that the agent will actually have some sort of inherent protections when it's trying to pull data out, generate data. It will then have to decide do I need to get approval for this or should the person asking for it really have access to this? Right.

Speaker 2: So if the content's not been, you know, previously tagged as confidential or containing PII or something.

Speaker 1: Right.

Speaker 2: You want the agent to recognize, oh gosh, there's PII here. You know, let me not proceed as I was planning to, but actually take a beat. Maybe ask a user for permission or follow some set of guidelines, but not just proceed mindlessly.

Speaker 1: I've heard this before from, from various tech leaders. You know, these AI agents are almost like our interns, right? Or, or any number of staff that they're going to work with us, right? They are, they are going to be helping us to iterate faster, to give greater efficiencies. But they need to have training, training the same training that an employee would have.

Speaker 2: Right?

Speaker 1: Security awareness training. Right. Privacy training. These AI agents, they need to have some sort of inherent guardrails, just like an employee would.

Speaker 2: What exactly do you mean by responsible AI at Walang duty?

Speaker 1: So responsible AI means that we have taken the artificial intelligence or LLM that's being used either internally or on a client project or to. To iterate for the client deliverable, if not the actual client deliverable. And we vetted it and vetting means privacy, security, regulatory compliance. It could even be data residency, right. That the, this LLM has to run, let's say in Germany and is not allowed to run in a different country because of processing and sub processing.

Speaker 1: Then also that we want to ensure that the model is being, while it's being vetted, we want to make sure that it's being run potentially in a secure environment, or maybe it's run locally, actually on a machine, in a hosted facility on, on premise. If it is coming directly from the manufacturer, whether it's OpenAI, Claude, IBM or others, we actually are even more careful about those models and we want to ensure that those go through the due diligence that every single application extension plugin goes through.

Speaker 1: It is no different. These models can be dangerous, right? And then there were so many unknowns that every human could not just check on. So we want to be more conservative about what we're doing, especially for anything that, that has direct access to our production data, confidential data, or especially highly

classified data. So we are very, very cautious to that regard.

Speaker 2: What business value shows up? When AI is governed correctly and executed correctly, we move beyond experimentation phase. Are you seeing tangible examples of business value being delivered yet? Do you anticipate them soon? And what does it look like?

Speaker 1: Tangibility is the key, right? What are we getting out of this? What is it beyond the R and D budget, right? For credits, what are we doing? It goes from show and tell and it goes right into outcomes. So the outcomes at long duty have actually been very good. Now, no, I'm not going to be speaking to 1x, 2x 10x processing or iteration, but what it's done is it's really transformed the way we think and our processes.

Speaker 1: The, the RFI RFP example is a great success that we were able to process, review and even respond or not respond to an RFP or RFI based upon our strengths as a company. The second example would be what the German and Serbia team are doing with AI, which is they're actually doing text image generation for clients and some of it is actually going into production. So those are the real outcomes. The of course IP and the indemnity are a big driver that are halting us from expanding this out beyond the clients who are willing to take some risks.

Speaker 1: So we are spending a lot of time looking at using AI internally for iteration, right? For design iteration. And then if we do take that to, to market for deliverable, then we still of course have humans that are leading it, humans in the loop, designers, strategy and all sorts of digital marketing staff that are still looking at the, at the data created or the analysis done by AI and making sure that it does respond in the way it's supposed to represent.

Speaker 1: Long duty and emphasis and we covered

Speaker 2: value and Driving outcomes. And for me, I think a lot about value, culture and experience is kind of three cornerstones of certainly, you know, my role in customer success. So I want to move on then to culture. And you used a term that I don't think I've heard before earlier, which you said human in the lead versus human in the loop. And so I'm sort of curious, what do you mean by that? I think I know, but I'd love to have you sort of explain what you mean by that.

Speaker 2: And then how do you apply that concept and what does it mean for your culture to say, you know, to your human employees? You're not, you're not just in the loop, but you're actually in the lead. And is that always true when you're leveraging AI? So just talk a little bit more about this human in the lead concept.

Speaker 1: Right. So for me, what that means is the origin of any tool that should be led by a human team and evaluated. Any tool needs to be continually, continually, continually vetted and it needs to be analyzed. Zero trust is very popular term in security architecture, privacy and compliance. However, we need to also extend that into our internal processes. For any time we're using a tool to make sure that it's actually doing what it's

supposed to be doing.

Speaker 1: Right. The worst thing we can do is, is is deploy some sort of tool and we put our trust in it and it doesn't give us the results and then we miss opportunities or just it hallucinates or just does the wrong thing based upon what we want it to do. So that is where the expectations need to meet the realities and we need to continually put that human in the lead to make sure that it is working and if it isn't, we'll change it.

Speaker 2: You know, looking forward, I think mostly of a forward looking question for you personally, what's your most controversial take on AI?

Speaker 1: Ethical use of AI? Not much discussion on that. As companies that are, that are, that are pushing it, are building data centers. Right. All these data centers are causing problems where they are built and they are not necessarily bringing value, the necessary value added to those local economies, of course, electricity generation effects to the planet. Those are some of the ethical things that are not being brought up very much. And that of course has to do with as the models get smarter, the AI credits for the smarter models will use even more power, water and effect to us as a society.

Speaker 2: Do you feel like there's solutions coming on that front or are you worried that there aren't?

Speaker 1: I don't know if there's solutions coming from that. I think that's going to be potentially a global problem because you can have of course, data centers that process the same LLM in a different part of the world that are going to offer lower cost per credit. So I really think it is a responsibility for governments and for companies to look at that cost of AI credits or which models are available and to whom they are available so we can all be responsible in the use of AI and not actually use it to to a negative effect to our global population.

Speaker 2: Well, that is a great way, I think, to wrap this conversation up. Jeff, I really appreciate your time, your partnership with us as a customer of Box and your insights here today. I hope that we this conversation was useful for all the folks who are watching or listening. And if you found it interesting, if you found it helpful and informative, please feel free to share it with your colleagues and I will see you on our next episode.

Speaker 2: Thank you.

Speaker 1: Thank you, John.

Speaker 2: Thanks for tuning into the AI first podcast, where we go beyond the buzz and into the real conversations shaping the future of work. If today's discussion helped you rethink how your organization can lead with AI, be sure to subscribe and share this episode with fellow tech leaders. Until next time, keep challenging assumptions, stay curious, and lead boldly into the AI first era.