

EPISODE 87**[INTRODUCTION]**

[0:00:06] ANNOUNCER: You are listening to 10,000 Swamp Leaders, leadership conversations that explore adapting and thriving in a complex world, with Rick Torseth and guests.

[INTERVIEW]

[0:00:20] RT: Hi, everybody. This is Rick Torseth and this is 10,000 Swamp Leaders. This is a home base for conversations with people who have made some decisions to use themselves to be of assistance, or what I would say aid the common good, both educationally and based on experience. Today is a return guest from 10,000 Swamp Leaders. He loved the swamp so much, he wanted to wade back in. Mick Yates is back with us. In a minute, I might bring him in here, but for those of you who are regular listeners will know that Mick and Martin Thomas were on an episode a about a year ago, I think Mick. I will post a link to that episode as well in the show notes, so that if you want to catch up on the first conversation, you can do that. First, Mr. Yates, welcome to the swamp again. How are you?

[0:01:03] MY: I'm doing fine, thank you, very well. A little bit warm over in the UK, but fine.

[0:01:07] RT: Good. All right, so we're going to talk about artificial intelligence here in this conversation. But before we do that, just share whatever you want people to know about you and why you're in this context a little bit and we're going to get more into the weeds on that, but what do you think they need to know about you to get going here?

[0:01:24] MY: The link back to the Martin Thomas conversation that we all had a little while back was leadership. We all went through the same program at Oxford and, and as you say, on consulting and coaching for change. One of the first things to know about me is I've been both a practitioner of leadership on the running business side, but also a consultant on leadership on the other side and a student on leadership. If I go back even before that, I've always been a bit of a data person. My first degree was mathematics and philosophy, which is basically logic. I've

been around computers in different ways for a long time. First ever programming was in Fortran on punch cards and IBM Big Iron in the 1960s, or something.

Perhaps more important to this conversation, I spent 10 years in the, what we then called the big data business with a company called Dunnhumby. They invented, particularly, a modern loyalty program where you use consumer data to fulfill their wishes and also send market into them, so and so. I'd always been using data in my previous career, but that was a big part of this. Bring it totally up to date, I'm a visiting professor at the University of Leeds in interdisciplinary ethics, which is a fancy way of saying ethics that has all kinds of angles to it and not just a particular subject area, where I occasionally teach on data and AI ethics. As part of that, I was a member of a group that was advising the EU before they introduced their AI legislation. That's the background that's relevant to this conversation, I think.

[0:02:58] RT: Okay, so let's pick it up there, because I know from time spent with you that you have been on the journey of understanding artificial intelligence, I would say, arguably long before most of the people that we know together in the [inaudible 0:03:12] alumni association. You're an early explorer of that world that was coming and now it's in everybody's lives. How did you get on that journey and bring us a little bit along that journey to roughly present day, if you can?

[0:03:28] MY: Yeah, okay. So, one of the key technologies behind artificial intelligence today is machine learning. But all machine learning isn't artificial intelligence. I mentioned that I was working with a company called Dunnhumby, both as a consultant and actually as an employee. This was, I guess I started 20 years ago. That whole business was built on machine learning. It was using machine learning to investigate customer information and get new insights from it. From a technological point of view, I've been around this stuff for, well, yeah, 20 years. I think my interest has just grown. As you've watched what's been going on, there's always been some new development that takes it to the next stage.

One of my other interests is photography. You might say, well, what's that got to do with anything? Well, what's interesting about today's AI, they're all extraordinarily good at reading images, because computer vision technology has really taken off in the last 10 years. I mean, it started back in World War II trying to identify targets for the military. Now it's detailed reading of

images, not just for what's in it, but perhaps what the meaning is, the ethics of it are. I think it's that convergence of different technologies that's pushed me along.

Where I am right now is actually motivated because of my philosophy background. I was at a biennial conference. Not biennial, it's every 20 years, or whatever that is. It was a conference which had – it was an interdisciplinary ethics conference at Leeds, and the subject of AI came up almost to a person every philosopher was saying, “No, it's impossible. AI cannot do philosophy.” I decided to go out and build myself on AI philosopher. Now, I haven't completely succeeded, but that's what I'm currently trying to do. It might sound a little silly, but if you think about AI, when you read the news, you see all the time about scientific breakthrough has been made, or a mathematical breakthrough has been made, or computers winning at the game Go, and so on.

When an AI is faced with a bounded measurable problem, a scientific problem, it can actually help human innovation, and that's now been quite well documented and proven. When it's dealing in less bounded subjects, conceptual inquiry, philosophical inquiry, scenario planning for the future of humanity, or whatever, it's less obvious how to use AI, because those are unbounded areas of interest and often, it's dealing with wicked problems and so on. Trying to get AI to be helpful in that space is really quite difficult. It's perhaps a long answer to your question, but it does perhaps explain this mixture of my desire to understand the technology, but at the same time, trying to make it useful and useful in a way that maybe other people aren't trying to do.

[0:06:20] RT: The word I'm wanting to use is translate, but that's probably not quite accurate, but just blending of AI and the philosophy. What have you learned so far specifically that you think is a contribution to this evolution that you're traveling?

[0:06:37] MY: I think that's a great question. I've learned it's very hard, basically. If you follow the technology of AI, you'll read about agents all the time. If you go a little deep on that, it will talk a lot about harnesses and it'll talk about human reinforcement learning and so on. It's all about how do you get an agent, an AI entity to do things in the way that you want it to be done? That's one of the biggest differences with the technology of AI with all the other great human innovations. It's not like the Internet. The Internet can't wake up on a Monday morning, post

something about you on its own and get you into trouble with The New York Times. Whereas, actually, an AI agent could do exactly that if it isn't under control.

This idea of agency is rather fundamental to what's going on right now with AI. To try to get the agent to do something which is unbounded in conceptual inquiry, as opposed to something that's bounded, find me a molecule that will solve this particular medical problem. That's really quite hard. But some of the same principles apply. There's a technique called multi-agent debate where you get agents to take certain positions, AI agents to take certain positions, some may be based on historical figures, some may be based on current social themes or whatever, and then debate those and then have adversarial agents to check if there's anything new, or interesting coming out of it.

My intent in this is not to try to find the right answer. The meaning of life is 42, according to the Hitchhiker's Guide to the Galaxy. That's not the point. It's to try to figure out how you can get AI to look at the different angles around an issue and give you some breadth on an issue that you might not be able to get on your own. Look at the tensions and disagreements and where they might be spaced for some kind of improvement in innovation. It's quite an open-ended quest, which makes it very hard. But at the same time, it also makes it really interesting. That's where I am.

[0:08:42] RT: Okay. I'm imagining somebody walking on a walk listening to our conversation and it's way behind you and the AI deal. Give a primer on agents, because that's a term that you're comfortable with, but I'm sure most people aren't there yet in understanding agents in AI.

[0:09:00] MY: That's a fair challenge. There are lots of different kinds of AI is the first thing to understand. We tend to bucket them in common discourse as AI, when what we really mean are chatbots. We really mean ChatGPT, or Claude, or Grok, or whatever. Now there's lots of other AI's doing all kinds of other things. But in common discourse, it tends to be these large language models, which are being talked about. A large language model is basically a prediction machine. It looks at all the corpus of human writing and texts and so on, and figures out what the next most likely word is in a conversation. It's a huge statistical word grinder, basically.

What that can do is if it's given the right kind of protocols, is it can take on a viewpoint about something. I could ask my chatbot, pretend you are a McKinsey consultant, a partner, and you're advising an oil company on doing something in Greenland, just to take something topical, okay? "What sort of research would you do and how would you go about it?" What you're doing is you're giving a definition of a project, or an area of interest to the LLM, and that becomes something that goes away and does. Then you can ask it to do something else. This morning, I was using it on my home inventory, because I was trying to figure out some of our artworks. I'd separately got Claude, which is one I tend to use. Okay, I want you to go away and find me all the references to those kind of artworks.

Instead of using ChatGPT as a kind of Google on steroids, you're giving it a whole raft of instructions to go and do something very specific. You can take those relatively simple examples and build it into something really quite important. You are the finance manager, finance director of my company. I want you to run my entire financial back office. I want you to prepare all my reports to the IRS. I want you to prepare all of my shareholder reports. That's your job. Go away and do that. Okay. Now, that wouldn't be a single agent that does that. What would happen in that situation is the LLM, the AI, would spawn lots of agents to go and do specific kinds of tasks, so sub-agents. Even on some of the simpler things that I've been talking about, that I'm doing, routinely, Claude would set 10, 20 different agents doing sub parts of the task reporting back, then trying to figure out where the answer is, or what's right, or what's useful

This idea of agency is very powerful, because it allows you to probe into particular areas. As I mentioned at the beginning, it's been proven that AI can help design new drugs, for example, with agents that are particularly good in certain areas of drug discovery, or looking at the literature, or whatever, understanding the ethics of drug testing. All of those are different agents doing different things in the system that actually delivers the drug. It's very powerful. But also, if it's not under control, it can do weird things. You've seen in the news recently, some of these high power LLMs escaping from their laboratory and their agents going off to do things, hacking Huggingface, or whatever.

What's doing the hacking is an agent of the LLM that's doing the hacking, because it's been tasked to go find out how to do that. An agent is this combination of technology and task and

harnessing that will get it to do certain things, which again, to repeat myself a little bit, it's super powerful, but it's also super scary.

[0:12:47] RT: Scary. Okay, so a lot of people are very, very much on the starting point of the path of understanding AI and how to use it in their life and using it, I would say, in a fashion you're describing, as opposed to a robust search engine. Let's assume that and then with your experience, how would you advise people to sharpen their direction around how they come to understand AI and make good use of it, rather than just bounce around pillar to post? What's the journey of learning that you think is some essential steps that need to be there in order for them to make progress?

[0:13:25] MY: I think the first and most important thing is to have some kind of purpose in what you're doing with the AI. If you're simply using it as a Google on steroids, that will be helpful to you, and we all do that. "Could you help me fix my connection between my TV and my new hi-fi unit or something? I don't quite know what's going wrong." AI is brilliant to that kind of stuff. But it would be more useful to say, "I've got this thing I'm trying to figure out." I am trying to figure out, for example, I've got a work problem. I'm trying to figure out if I should be looking at customer segmentation in my work a different way. I've got a group I'm working with on photography and I'm trying to help them improve their output, improve their creativity. Can I take some of their photographs, can I analyze it, can I suggest improvements? I would say that starting with something that is important to you that you really want to try to understand is one of the most critical first steps.

One of the great things about AI is its breadth. I mentioned that I'm a part-time philosopher. Well, it's interesting in the world of philosophy, like a lot of academic areas, it's very specialized. You don't talk to a moral philosopher too much about aesthetics. That's the way the world is. We've all got more specialized as time goes on. The great thing about AI is it will offer you different kinds of perspectives. I'm particularly interested in the work of David Hume. He was one of the great empiricists, one of the most important philosophers in history in many ways, and he was always having this debate with Immanuel Kant about knowledge and intuition and so on. AI can help me understand that debate, but it can also give me perspectives on it that I wouldn't have thought of.

If I had Immanuel, David Hume, and Buddha and Sun Tzu in the same room talking about empiricism, what would happen? How would that work? What would happen if I took somebody you work with, an adaptive leadership expert, like Ronald Heifetz, how would he get on in a conversation with Machiavelli? When the AI helps you with those conversations, it isn't repeating everything from the textbooks. What it's doing and the way that I use it anyway, it's creating patterns of thought, the kinds of things that Machiavelli might have thought about leadership, or Buddha might have thought about empiricism, or whatever. Using those patterns, in debate against patterns coming from David Hume, or Immanuel Kant.

That's perhaps a more complicated example, but that leads you a little bit into what I'm trying to do with philosophy. It's the same point I was trying to make. You've got this agency. You build a question that you really want to get answered. You build some methodological guidelines, if you can. You give it areas of interest, areas that are in bounds, areas of out of bounds, and off you go. I think the more you work with it, frankly, the more you'll find it useful.

Anyway, this isn't the same what I'm describing as prompt engineering. It is in one way, and the prompt engineers will say, "Well, I'm a media expert. Give me a plan of advertising over the next 12 hours." That's a little bit. What I'm saying is a bit more than that. It's like considering the possibilities inside the LLM of its breadth, depth, different things it could cover, how it could take different disciplines and smash them together. I'm very interested in interdisciplinary work. See what it can come up with. It leads naturally to things like in the world that you and I have it, like scenario planning. When you're looking at lots and lots of different influences, how can you get the pattern matching of an LLM with different agents, understanding different parts of this scenario situation you're trying to work your way through. How can it be helpful to that? You start with an interest and you gradually build it up into something that, yeah, this is something big that I can get my teeth into.

[0:17:21] RT: Okay. At the beginning of our conversation, you're talking about you also have a very strong exploration going on around leadership and you have for a long time, I'm thinking too about, and you have been a leader in large groups and large global organizations around the world. From a leadership and an authority position in an organization like that, now you've got AI running around in your organization and people sitting at their workplaces doing who knows what with AI. When you put that hat on, if you were back in those roles today with what

we're dealing with, what concerns and what do you think as a leader of the organization, you need to be keeping an eye on to manage that activity?

[0:18:02] MY: Okay. I think a couple of things, two buckets, a bucket which doesn't change very much and a bucket which does change. In the bucket that doesn't change very much, make sure you know what your strategy is. Let's just start from the beginning, okay? We're both old enough to remember the dot-com era, right? Everybody's rushing around saying, "I need to be on the Internet." People building websites that didn't do very much, didn't really help their business. Most of those things went wrong, because there was no real strategic thinking behind the intent of this site, how it was going to work for customers and so on. How does it fit into your business model?

Then, of course, you had the geniuses, like Amazon that came out and completely changed the game by saying, "Oh, well, we can do this with this technology." They really understood this technology and it took it in a different direction. The first thing is be really clear strategically what it is you're trying to do with the business and try to figure out where technology can be helpful to you. I don't think that's changed. That doesn't change whether you're talking about the Internet, big data, or AI. What does change, I think, however, is that you really do need to understand this technology as a leader. I don't think you can outsource this. I don't think you can send this away. Let your IT department figure out how to use AI. If you go back to everything I've been saying about the agency of AI, what the agents might actually do that can affect your business, that isn't just an IT issue. That's a whole of business set of issues.

I think the CEO needs to go to great lengths to try to do a better than average job of understanding what the technology can do for you, what its pitfalls are, and so on. Of course, get some people that really understand it. Of course, get the experts. I do think you've got to do a better line. If I was still running a big business, I would be really trying to understand how it works, what's in it, what are people saying about it, what are the risks and how can I work it. That's not something I want to outsource.

[0:19:57] RT: I would imagine, in order to get quite competent in a deep understanding of what it is, you actually have to carve out serious, consistent time in your agenda, because this is a big

deal and you're starting from scratch. And so, the challenge of how do you get it into the space of time, you've got to do all the things you've got to do in that position is its own challenge.

[0:20:17] MY: Yeah. But that's just like business as usual, isn't it?

[0:20:19] RT: Yeah. Yeah.

[0:20:20] MY: Any CEO, or senior business leader is going to have a ton of different things pressing on them at any one point in time, so this is just another one. It's just this one happens to be really quite important.

[0:20:33] RT: All right, so just recently and I'm sure you either read it, or watch the video of the interview that The Economist had with Elon Musk a couple weeks back. I watched it and he took, I thought it was interesting, he had two different views. He had this Nirvana description of the future where there's going to be so much abundance that people won't have to do anything, except just absorb the abundance and live life, and then a notion that he can't sleep at night, because sometimes he's scared to death. Now this is Elon Musk and let's take a bit with a grain of salt here. But when you think about the future, what are the things that pop up that you think this is good for the future if we can move the needle on this and this is to be wary of because this could be a problem?

[0:21:16] MY: Okay, so I'm not going to get into a head-to-head fight with Elon Musk, though that could be quite fun. Maybe not on this occasion. I think that if you think of AI as a replacement for human beings, which is a little bit how Elon was talking about that, that's going to lead you down a very dark path. If you think about AI as a partner, an assistant, a useful tool with unusual agentic capabilities that can be helpful, I think that can take you down a brighter future. I use the term orchestrated intelligence to think about that. When I've been using AI myself, I've never been trying to get it to replace me. I mean, it might replace me on doing an Excel spreadsheet. It's much better doing it so formally than I'll ever be, so that's great. But it's not at this point going to outthink me and it's not going to replace me as a consultant, or a philosopher, or whatever.

I think, if you think of it as a partner, something you've got to work with, and I think it's quite positive. But I also think you need to look at the technology. Humanity has had a lot of different technologies and every technology we've had since fire in the wheel, it always had dual purpose, right? It's either going to do a lot of good for you, it's going to kill you, right? Nuclear energy can do great things. It can kill you. The Internet per se probably can't kill you, but actually in some of the stuff that circulates around the Internet, it can certainly not do your mental health that good. It has this dual purpose. The same is true with AI. It has the power of great good and the power of great bad. I think it's closer in this regard to nuclear energy than almost anything else. Nuclear energy is quite regulated. You do not let your average CEO, or your average company run around building random nuclear weapons. That doesn't happen. Even the massively deregulated United States of America doesn't let that happen, right?

We've got AI leaders now almost unanimously saying, "Guys, we need to get some kind of regulation in the system." Now some of that's a bit self-serving. We haven't got time to cover that, but it's real because of this power to do harm and as well as do good. At the moment, if you look at it, AI is doing good for a relatively small number of people who are making a huge amount of money out of it. It's upsetting people with data centers. It's upsetting people fearing they're going to be out of a job and so on and so forth.

There needs to be a serious conversation and not the kind of Elon Musk conversation, frankly. I don't think most of that conversation is particularly serious. I think there's got to be a serious conversation about how can you create some kind of structures here to work out how everybody can use the good bits, the bad bits are kept under control and everybody can at least get something out of it. When electricity was invented, it wasn't just sent to the rich people. It was to start with, but now everybody's got it. Well, not everybody, but you know what I mean. How do we do that with AI? A lot of that is through regulation, social discourse and so on. I think that is very important for this. If I was sitting opposite Elon, that's where I would be taking the conversation. He'd be hard pressed to argue with me on that frankly, because even the libertarian would not want to have random companies creating nuclear weapons.

[0:24:55] RT: I'm going to see if I can get Elon and you on another podcast and we can go for that.

[0:24:59] MY: Yeah. Good luck with that.

[0:25:00] RT: Yeah. You spring two questions though, one that I thought of in advance and then this one here, which is given all of what you just said, what are, as you view it right now, this is obviously fluid, what are some sweet spots for human beings, people working that AI cannot at least at this point address, finding harmony and value in work based on doing something meaningful and having part of it being seeded by the fact that I'm a human being and I've got something that AI just can't bring to the party?

[0:25:33] MY: Right. I think there's two aspects of that. It's yet to be proven that the current raft of AIs, because they are disembodied from the world, we can come back to that in a minute, if you like. It's yet to be proven that they can truly create our literature, break through thinking even in the sciences in the way that humans can. I think it's probably true to say that an LLM, which is bounded in working out words in a data center somewhere in the middle of nowhere is not connected to the world, and so isn't going to have those experiences. It's not going to be able to – In Europe, we were all enjoying the eclipse yesterday to one degree or another. It was an interesting thing that happened. We were in a local pub and people were talking to each other about the eclipse and how does it happen and how can you see it safely and all the rest of it. It was a very human thing. We were grounded in the world.

There are attempts to do that with what's called world models to ground AIs in the world. At the moment, that grounding in the world gives a lot of free space for human creative thinking, strategically, intellectual thinking, games, pure creative thinking, and so on. We have this conversation about music and AI. That's going to keep coming up and going, and really, that's fine, but it's not going to destroy music. It's just going to be another thing people are going to have to work with, the same with the art. It's something that people have to work with. I think that's one.

On the other hand, actually, it's quite banal. I mentioned, now I tend to use AI just when I'm trying to do simple spreadsheets. They happen to be really good at doing things that I'm not very good at. They are really good, and they're not impatient like I am. They will change a spreadsheet 50 times. I never once say to me, "Mick, why didn't you get it right the first time?"

They're not going to do that. "Oh, yes. Okay. Well, we can change it around like this, or we can do this." Yeah, it's fine. Simply using it in your everyday life is actually quite a cool thing to do.

[0:27:38] RT: It's not stopping and having to go out and get some fresh air, because its brain is not working so well right now.

[0:27:42] MY: Yeah. It's still sitting there. "Well, what do what to do now? What can I help you with?"

[0:27:46] RT: My plan was to turn a little bit to your leadership background here and I'm going to do that, but I don't want to lose sight of the disembodied element that you raise. Speak to that, because that's new to me and I don't know what the question is I should ask. What is it about that that's relevant for people to have a handle on?

[0:28:02] MY: Okay, so a couple of pieces to this and forgive me for a little bit of philosophy on this. You'll see in actually, in regular newspapers, is the AI conscious? You see this all the time. Is the AI conscious? Well, and you've even got people like Jeffrey Hinton, one of the godfathers of AI thinks it might be. Anthropic recently came up with some research, very well-founded research, actually, talked about these spaces inside the LLM operations that didn't know existed, which is where the LLM had set aside some space, so it could do certain kinds of calculations to get its act together, basically. I think they were called J places. All of those things build on this mechanistic idea of the brain, that the brain is just a big computer with lots of neurons, lots of connections and so on. If we can do that with computers, eventually they'll become conscious. Well, the problem with that is nobody actually understands what consciousness is in a human being. Never mind how to wire it into a machine, right?

[0:29:06] RT: Right.

[0:29:08] MY: All that stuff about is it conscious and all the rest of that, it can maybe fake it, but it can't actually do it in its current configuration. If you look at the literature around consciousness, one of the things that I always come back to is this idea of being in the world. You look at a human baby, my wife, her and I are blessed, we have 11 grandchildren, and we've watched not only around six kids grow up and their partners, but also 11 grandchildren. They

didn't learn just by reading books, which is what LLMs have done. They learned by falling over and hitting their head. They learned by getting stung by a wasp, thinking it was a nice thing to play with and realizing it probably wasn't. They put their head in the water and find they can't breathe and they have to get out of there in hurry, right? They learn by experience.

In one of the next big areas of exploration in the world of AI is going to be what's called world models, where you create models of the world, so that the AI can experience that. That's not the same as a robot. A robot, as they're currently configured, are meant to do predetermined moves, even these fancy Chinese robots, they've got predetermined moves. Yes, they've got some latitude, and so on and so forth. They know what kind of kicks to make when they're fighting and so on. But they're not out there to understand what it smells like to be in a sawdust rink. That's not what they're for. They're there to just have a fight.

The idea of a world model is to give the world experience of a human and marry that with the idea of the AI's understanding and see what happens. That is going to be the next big thing. There are already people doing that. One of the biggest is actually being put together by Yan Lukan, who is one of the other grandfathers of AI. He was for a while at Meta, but now he's got a new startup based in Paris, London, and San Francisco, I think. Basically, a French company. That's where it's going. This disembodiment and fixing it, that's a huge area. That's not going to get resolved in the next couple of years. Does that make any sense? Is that going to your question? Or is that moving off that a bit?

[0:31:25] RT: Well, I think it's both. You're way down the road from where I am, so I'm swimming upstream to understand. For the broad end of listeners, they're going to pick up on this. I will say also for our listeners, and you and I can work this out offline, you probably have some other resources that you can send me links to that people could dig into after they've listed you and we'll put those in the show notes, so that people have a next step in the journey, both your own writing and others.

Let's just shift a little bit, because another big part of who you are is your work as a leader. You developed several years ago a framework called 4Es, which stands for Envision, Enable, Empower, and Energize. I was familiar with this from way back when, but I got re-familiarized with it when I was preparing for our conversation here. I think this is an important piece in the

world of AI and where people who are having to use themselves to lead could benefit from some actual robust understanding about how they might do that and use their capacities of human beings. I think your model gets to that if I understand it well enough. Speak to that a little bit, because that's an important thing you've added to the legacy here, but it fits in here somewhere.

[0:32:33] MY: Yeah. Thanks for bringing this up, actually, because I think it does. The genesis of this goes back to when I was working at Procter & Gamble, an assignment way back when in Cincinnati in the 80s. They were, like a lot of corporations, starting to kick around with leadership models and so on. They were looking at things and they were using Es. Not these Es, but they were using Es as a metaphor. I did some work on that. Then I was using that practically, and when I went on to my next assignments in J&J, and so on. Then, when I did the course at Oxford that you and I both did Oxford and HSA Consulting, because coding for change, I decided to make this my master's project. And so, I researched it properly. I did some case studies.

At that time, I was on the board of Save the Children in the US. I did a case study looking up leadership in NGOs and using Save the Children as an example. They allowed me to do that. The idea of the Es is that it really starts from Warren Bennett. Warren, who I was privileged to know, was probably the first person to say, leadership was a process. It wasn't like a character trait. It was a process. I tried to get inside that process. This idea I've been visioning the strategy first and enabling the tools, a bit like we've been saying on AI. What's the strategy? What are the tools? Empowering the organization to get on with that and figuring out what the right structure is to do that.

Then you as a leader energizing the whole system, walking the talk, making sure you're ready to course correct, and so on. That's how that model was constructed. It was meant to say, well, dissect the process and take it in new steps. I think the thing that I was most pleased with in the model, I broke it into two axes, operational and organizational. Because you can think about the four things I've just said, what's the purpose of my organization? What are the tools for my organization? And so on. But then, if I'm thinking about the task at hand, what am I trying to do? What's the output of my business? That's a different axis. Putting the Es on those different axes, I think is really quite helpful.

When I did this, it wasn't meant to say that other models were bad. That's just not the way my mind works. It was to say, this is a way to think about the process of leadership and overlaying it on whatever else you might be doing in the leadership space. I think it can be applied to AI, as I've said, I think it's the same steps. What are we trying to get done? What's the technology? How do I get my organization lined up around, and so on? It's not rocket science. I just think it's a useful frame to hang these things up. I think it definitely can come into its own in the AI conversation.

[0:35:19] RT: Okay. I put this in my questions and I'm going to admit that I'm not quite sure where I found this inside the stuff you sent me, but you will know. This may or may not be the right place in a conversation to this, but it's the right place on my call sheet here. I want to know, you raise a question, how do we maintain human agency and motivation while leveraging machine intelligence? Now, I think we've touched on this a little bit, but I just want to sharpen it if we haven't. What is your answer to your own question at this point?

[0:35:47] MY: I think that's a great thing to pick out. I think it's an unanswered question. Because if you listen, to go back to Elon Musk, you listen to what he's saying, we won't have any money, eight machines will do everything. What are we going to do? I mean, that's a pretty dreadful looking future to me.

[0:36:02] RT: Me too.

[0:36:03] MY: That doesn't have too much agency in it with a human being. I think if you take it as a tool that we can use to help us, again, what I have described as this orchestrated intelligence in some kind of partnership, helping us further our thinking, then I think it can be – That's a very useful way to think about this. I don't know if you saw on one of the groups that we we're in, I was doing a little bit of research on Go. What was quite interesting was the way that after – So, back a step. Some years ago, the AlphaGo machine beat the World Go Grandmaster, Lee Sedol, by making some moves that nobody had even expected were possible. It literally blew everybody away. They were almost at awe of these moves.

If you now look at a systematic view of all the smart Go moves that humans have made over the decades, there's been a whole bunch of new ones been made since AlphaGo did this. It's

inspired Go players to do more. If you take AI in that way, and you say, "This could help me do things, and therefore, I can do more," I would see that's part of the answer to the question. If you take AI as a way to eliminate human labor, which is frankly, what some people are trying to do, then I think that would be a negative answer to that question. Does that start to get close to what you wanted to explore here?

[0:37:33] RT: Yeah. It was your question, and I was wondering, A, digging deeper into the question which you've explained, and also where you are so far with perhaps, arriving at a conclusion. I think you've covered that, too, a little bit.

[0:37:45] MY: Yeah. I mean, I've got a conclusion in terms of where I want it to be. Now, can I guarantee that's where it's all going to go? Absolutely. I mean, right now we've got a relatively small number of hugely influential and very wealthy people driving this stuff. I think that's got to change.

[0:38:02] RT: Yeah. Okay, I want to come back to that one too, because it's down here. You have a thing called the – and you just mentioned this not in the specific, but let's sharpen it. You have something you call The Studio of Orchestrated Intelligence. Tell people what that is, because that's a thing, people, as I understand it, can engage in as if they so choose. What is this?

[0:38:22] MY: Yeah. It's in beta test right now. It goes back to what I said at the beginning about trying to build myself a set of tools to investigate concepts, rather than build new molecules. I have created what I call The Studio, which is a set of protocols, code, predefined agents, predefined methodologies with quality controls built in it, so that we can manage hallucinations, and so on. We can understand the pedigree of what's being said. We can understand whether it's got any empirical basis, or whether it's completely vacuous idea. We can see if it's got some innovative possibilities, or whether it's frankly, something that's a bit of a dead end. I've built something with these protocols and quality controls, which I call The Studio, which I can interrogate to ask questions. They could be simple questions. For example, we've used it to advise small business owners, simple questions. Should I give my best friend better pricing than I give my best customers? Okay. It's quite interesting to see how an AI deals with that.

It isn't just a prompt. It's actually using a set of protocols and methodologies and different approaches to ethics. In that simple question, you've got utilitarian ethics, you've got the golden rule ethics, you've got all kinds of different ethical questions in it. The answer to that question might be different in a Western context, or an Asian context. The tool is designed to explore those kinds of things. I've been writing papers on it with other folks. Like I said, we currently have it in beta test.

Related to it, there is a thing that I call The Gallery. Again, that goes back to what I said about computer vision. It is actually extraordinary, how even an LLM frankly can interrogate an image. My other life, I'm a photographer, I teach photography, I do things in that space. I'm interested in the ethical side of photography, what's a good image, what should be in front of the people, how do people respond to it, and so on. Actually, an AI can help in that interrogation. The tool, I call that The Gallery, that's another tool in this. Then it's all bound together. The old statement is always true, garbage in, garbage out. I have a knowledge management system to try to keep abreast of what's new in different areas of interest for me; technology, photography, whatever, which I call The Library. None of these names are very clever by the way, The Studio, The Gallery, and The Library. But hey, it works. I've been trying to build this. Rather than just talk about this, I've literally been trying to put my money around my mouth is and building tools to be helpful.

[0:41:14] RT: These could be things that people eventually would access themselves then and use? Is that how you're thinking of it?

[0:41:19] MY: Yeah. I'm trying to figure out how to do that right now. I've got it in beta test. I've got a few people from our mutual group who are using it. I've got some academics been using it, small businesses using it to try to make it robust, make it work. Again, I'm not looking for some really big breakthrough. This is a massive new idea, which is going to build you a trillion-dollar business. I'm trying to figure out how to use AI to help people explore concepts, questions, wicked problems, and give them clues. That's really what it's all about. I can send a link to where it is and people can have a look, at least learn about it on my website and if people are interested and beta test it, then get in touch with me. Yeah.

[0:42:01] RT: Great. That'd be ideal. All right, so that's a nice lead. Among all the things you've been up to, you've been a visiting professor at Leeds University for a long time, I think 15, 16 years. Therefore, you know a bit about teaching in a classroom, and AI has all sorts of ways in which you can do that, but you have a massive amount of actual experience. If you had four hours and a room full of people who are somewhat aware of AI and they use it like a robust search engine, how would you construct, or what would be the elements that you would want to cover in the four-hour period that would advance their understanding and their ability to use AI as an instructor?

[0:42:44] MY: Okay. Well, I think, first of all, I wouldn't do it on my own. I think it'd be really important to have practitioners show what they're up to, because there's an inexhaustible amount of activities out there on AI. Looking at different ways to use AI, I think is very important. An instructor can cover some of that, but there'd be no substitute. I mean, a couple of other people showing what they're actually doing. That's one thing I'd definitely do.

Second thing, I would talk about the connection, which I think we've covered pretty well in this podcast, about the leadership role and the technology role. Because there's not much point if you're the kind of people that you and I would be working with, business leaders, consultants or whatever. There's not much point them doing this theoretically. They need to be put in a context of what they're trying to do with their business, whether enterprise, or organization, social network, whatever they're trying to do. I think I would absolutely try to connect it into the leadership role in AI. But then, I absolutely would also cover two things. I would cover some of the technology. I hope I've explained some of it a little bit today, but I think there's a lot of misunderstanding about technology and sometimes I get it wrong too, frankly, because it's moving so quickly.

I do think understanding a little bit about the technology, you know, what is an agent? How does that work? I think that's really important. I don't think you can just ignore it. You can't just say, "Well, that's just like HTML was for the Internet. That's just our website is built." Well, no. An agent is something that does something that can impact you, so you better understand what it is and how you use it. I would also try to address, although it's four hours, so we're running out of time here, I would also try to address a little bit about this combination of the ethical challenges of AI, regulatory challenges of AI, and the future of work. Because I don't think – the world's in a

funny place in a regulatory space. The world's in a funny place. The EU has done, I think, a commendable job. I know I'm biased, but a commendable job of trying to figure out how to think about AI from a personal, individual citizen point of view, and from a risk management point of view.

The Chinese government have got more regulations than anybody else, but that's largely about controlling the technology, having state control over it, which is probably necessary, but it doesn't deal with some of the other issues, obviously, because it's driven from a political perspective. Right now, the US is a desert when it comes to regulatory. It's a desert, because I don't want to understand it, or they're making too much money on it, and so on. But as I said, the people that are leading frontier AI models are starting to say, "Guys, we need to think seriously about this regulatory thing."

We've got these three ways of looking at it. There's no such thing as nuclear fusion, or nuclear fission, which is Chinese, or European, or American. Nuclear fusion or fission is nuclear fusion or fission globally, full stop. Right? AI is the same, full stop. Although it's got political implications, it's still going to be a technology that's everywhere. I think getting to grips with regulatory, having a bit of a conversation about that would absolutely be part of whatever I do in four hours.

[0:45:59] RT: Okay. Well, we'll probably stretch it to six or seven hours just to give you more elbow room there.

[0:46:05] MY: Probably.

[0:46:06] RT: All right. While I was preparing, this idea came sideways across my head and I thought, well, what the heck? I've never done this before. I decided, I use Claude, and the first question I asked Claude is, can you find Mickey H in his websites? That was a low-hanging fruit. It found it. I asked Claude what I should ask you and to give it to me in some questions. It did that and I called them down to three. I think it gave me 10, or 11. I don't know. Here's question one, and this is about the co-work transcript question. This is your context that you've set. Claude said, "Mick use Claude to help write the essay that concludes AI can't produce genuine novelty." I don't know if that's true or not, but let's go with that. Then it says, "Ask Mick, in this

specific Claude session, was there a moment of the tool produced a connection he didn't see? If so, does that complicate his own conclusion? Or does he have an answer that survives it? This is checkable, not abstract. He has the receipts.”

[0:47:11] MY: Yeah, that's typical Claude. The answer is, yeah, there's definitely times when I've seen things that would classify as innovative thinking. Some of the quality control mechanisms that I use on the studio that we talked about have been invented by Claude, not by me. Do they classify as world shaking inventions? Probably not. But they absolutely classify as innovation. To Claude's question, yes, they would survive, because I have documentation. I have a set of memory artifacts. You should put this back to Claude, which captures the most important findings of most of my important sessions over the last 12 months. In fact, my Claude is pushing me to do a big interrogation on our data set to look backwards. Tell your Claude that.

Yes, there are definitely things that have come along and say, “Well, I didn't think of that. That's clever. That's new.” It also happens with synthesis, being able to put something from a domain of knowledge that I really don't know much about that Claude would help me with, or my sister would help me with. That happens all the time. That's checkable, because I can then go back and say, “Let me find out if that's really true, whether this is being made up.” Anyway, there you go.

[0:48:31] RT: Okay. One more question from Claude. This is the headline. Is he the thing he says the discipline requires? That's the question. He argues that “becoming has to be protected by a person, or a practice operating outside the archives pulled towards convergence. He also built a very well-organized archive, a wiki, The Studio that he feeds himself into. Ask him straight” I'm going to ask you straight. I'm just going to read the question here. Is he currently that outside voice, or has he built the very system that would eventually absorb him?

[0:49:05] MY: Yeah, that's a got you from Claude. There is a massive tendency when you're building a knowledge management system without AI for you to build that in your own image. That is how it works. I mean, you're a consultant, Rick. You know you've got all kinds of stuff on adaptive leadership, which is your thing. That is absolutely the case. The only thing I would say to Claude, your Claude, is that he may be able to find reference to what we call rule delta one, which is that the apparatus cannot mark its own homework. My system, my studio system, calls

rule delta one all the time, because the idea that you can have an AI system creating these wonderful reports, McKinsey style reports, and then rating them and saying how good they are. It's not like a self-enclosed thing. IT'S being absorbed. It's exactly what your Claude is saying. It's all being absorbed into one system. It is an absolute got you question. Your Claude is right. But there are ways to try to deal with that. My rule delta one on some other things. how I try not to become the Borg.

[0:50:22] RT: Well done, sir. All right, so we're coming down to the end here. What have I not asked you that you want to throw in here before we complete?

[0:50:30] MY: To be honest, you've covered an awful lot of ground. I hope you and I hope your listeners have found this with some interest. It's not something you haven't covered, but I think it's perhaps my biggest suggestion out of all this is for people, please, please, please, educate yourselves on this stuff. I think it's really, really important. It's quite possible that 75% of what I said today is complete rubbish, right? I'd rather have somebody find out for themselves and tell me what's right or wrong, than just assume what they're reading in the newspapers, or assume what they're listening to in a podcast is the right solution here. I think, please, please, please, educate yourselves. Find out what you can. Play with it. It's got infinite patience more than you've got. It's got more data than you've got. It can be super, super helpful to you. It will probably annoy the heck out of you, but please educate yourselves.

[0:51:30] RT: Right. Nice ending. Mick Yates, thank you very much for making time to help people understand a bit more about AI than they start this conversation with. I know you've got them further down the road. Thanks so much.

[0:51:42] MY: Thank you. Thanks for having me.

[END OF INTERVIEW]

[0:51:45] ANNOUNCER: Thank you for listening to 10,000 Swamp Leaders with Rick Torseth. Please, take this moment and hit subscribe to follow more leadership swamp conversations.

[END]