

PODCAST TRANSCRIPT

Talking Research: Why AI heralds the next Industrial Revolution

First published: 26 August 2026

Katie Boyce: Hello. I'm Katie Boyce from Walter Scott. Welcome to a special episode of Talking Research. Taken from our recent research conference here in Edinburgh, we thought you'd enjoy this conversation between Microsoft AI's Michael Bhaskar and investment manager Tom Miedema.

Before we get started, a reminder that this podcast is intended for investment professionals only and should not be construed as investment advice or a recommendation. Any stock examples discussed are provided in the context of the theme being explored, and the views expressed are those of the presenters at the time of recording.

Tom Miedema: Good afternoon. For those who haven't met me, my name is Tom Miedema, part of the research team at Walter Scott, investment manager. And today, it's my pleasure to introduce Michael Bhaskar.

The topic here is AI and Michael is a thoughtful and wide-ranging voice on AI technology, and many other things: he is an author of four different books, he's a researcher, a strategist, an entrepreneur. For a lot of his career he's had the privilege of watching and being part of the evolution of AI.

His current role is leading Microsoft AI's Futures team, a team that's responsible for strategy planning, for thought leadership around many things to do with AI, and a lot of the papers you can read on the website. Before that, he was Writer in Residence at Google DeepMind. 'Writer in Residence' is just a brilliant job title.

I first came across Michael's work reading the book, *The Coming Wave*, and he's probably best known for co-writing that book with Mustafa Suleyman, who is today the head of Microsoft Superintelligence. For me, I think this book is now required syllabus for getting to know AI. I first read it last year and would highly recommend it to anyone.

The plan is to have a really broad-ranging conversation about AI and touch on many different topics. So, with that, Michael, welcome.

Michael Bhaskar: Thank you for having me. Great to be here.

TM: I was going to start with the technology. You've been watching this space for a very long time now,

and it feels to us like a lot has changed over the last six to twelve months; GenTech AI, big step up in model capability. But I'd love you to set the context for someone who's really been in and about it for all that period; how do you see things?

MB: I think there are two views. There's the slightly more long-term view, and then there's the near-term view of what's been happening over the last twelve to six months; this time last year, everyone thought we were going to be hitting a wall. They thought that progress was going to be stopping. There was a lot of talk about the capabilities starting to tail off.

Spoiler alert, that didn't happen; actually the opposite happened. Capabilities started improving at a pace that was even beyond what we saw. But I'll come back to that because the recent progress is just built on the back of a very long-term set of trends. This modern revolution in deep learning, which underlies AI, has been going on since about 2010. Something that our team calculated is that since 2010, the amount of available compute for frontier training of AI models has gone up one trillion-fold around the world. There's just a trillion times more compute available to make AI.

If you believe in something called the scaling laws, then this basically explains everything. The scaling laws say the more compute you have, the better the AI that you get. And what's remarkable is, that has held over all that time. In 2010, people didn't even know about the scaling laws. It was only nine or ten years later that people observed that that relationship was there.

But the scaling laws are still going on. The compute is increasing. What's underlying that increase in compute is GPUs. We went from CPU systems to graphics processing units, much better for AI. For the kind of maths that AI needs, GPUs are perfect. People tend to think that's the hardware story. That's not the hardware story. The hardware story is also about memory. There have been all kinds of revolutions in memory, which means that you can put information into the GPUs faster. So, you've got better GPUs. You've got more memory. Then you can wire the data centres together much, much better. You can interlink all of these GPUs. I'm sure the photonics that we just heard about is actually about interconnecting different chips in a data centre.

When a model called AlexNet came out in 2012, that was pretty much a frontier model at the time. It was what established deep learning as a really breakthrough technology. It did image recognition. That used two GPUs. Now it's common for data centres to have 100,000 GPUs all interlinked, all working together on a single training run, and each one of those GPUs is far, far better than those two GPUs. This is just this wild story of multiple different technological exponentials, and you can really laser into each one of these technologies and see countless improvements just delivering more compute.

And then, on top of that, you have algorithmic breakthroughs that let you do more, but it's really that macro picture that has sustained the progress over time.

And that brings you to the last six to twelve months. What's actually happened? The big thing is reasoning models. What you have there is the reintroduction of a technique called reinforcement learning into large language models. You're all familiar with large language models. You're probably all familiar with the technique that they use, which is next-token prediction. They're basically just guessing the next word all the time. That is one way of doing AI.

Another way of doing AI is reinforcement learning. The AI tries something, if it gets it right, it gets rewarded. If it gets it wrong, it's told, "No, that's not right," and over time it comes to understand what's right or wrong. To build these reasoning models, they introduced that RL process. They start saying, "You're getting it right. You're getting it wrong." And this has been massively powerful and has really unlocked the next layer of capability. And probably starting in November with the launch of an upgrade to the Claude Opus model, we started to see just incredible gains from what the AI could do.

And really everything switched again in that moment, probably the biggest breakthrough since ChatGPT launched in 2022. What it basically means is that workflows that even a year ago would have been impossible for AI to handle, it can now really handle. What it means is that AI has, in my view, gone from being an absolutely fascinating thing to a business-critical thing. And that's the long-term picture of this insane rise in compute now unlocking these ultra-capable new generation of models, specifically coding models.

TM: Do you want to touch on the agentic shift as well? The first time we spoke, you were very excited about agents and Claude, et cetera. So, just touch on how agents have evolved as part of this as well.

MB: All of these models that are doing so well can create and run agents. I'm sure everyone in this room will have heard about agents, but simply put, agentic AI is when it stops just saying things and starts doing things. It's not just giving you advice on producing

the PowerPoint, it's producing the PowerPoint. It's not just creating a shopping list for you, it's actually buying it all for you. It's doing whole projects. It's executing tasks on behalf of businesses.

People started talking about agents probably in 2024. And I'd go to conferences and people would say, "Ah, the agentic era is here." And actually, it wasn't. The agentic era probably arrived about six months ago. But now very quickly, everyone on my team, everyone at Microsoft AI - and I mean everyone - is spinning up and running multiple teams of agents to perform things. You will specify a set of tasks and then say, "Right, create a team of agents that'll go off and do this portion, create a team of agents that'll do this portion."

There's been a huge breakthrough in the interface for agents with each other and with things on the web, protocols that allow them to interact with tools. It really is a time when all of this talk about agents is starting to be realised.

There's a really important benchmark, probably the most important benchmark in AI, and that measures how long agents can go off and autonomously perform a task. And it's probably now getting up towards a day for the very best agents on coding tasks. Now consider that a few years ago, it was a couple of minutes; that is an agent could work as well as a good or well-qualified human being for a couple of minutes. It can now go off and do it for roughly a day of a senior coder; and at AGI level, it's probably about a month. So, we've still got a little way to go for that. But what was interesting on this time horizon task chart, is that there was a smooth upward curve, and we're now going beyond where we thought we would be on it.

TM: It's really quite a change. We've had that internally as well. If I can go into Microsoft for a bit and talk about what's happening there specifically, I think it'd be fair to say that many people here will be users of Copilot. Microsoft's incredibly dominant enterprise software with the Office Suite, et cetera. When you're looking at the comparison between Copilot and Claude today, or Claude Cowork or some of these different harnesses, it's improving, but it's still, I think, far from up there with the best.

Internally, it must be an incredible challenge because, obviously, it's Microsoft's market to win. So maybe just speak about internal efforts or how you view that competition.

MB: I'd start by saying, I sit in the internal models team. In terms of the Frontier lab, we're a little removed from the product. There is an extraordinary amount of work going on to make sure that all of Microsoft's products are going to be, not just competitive with the best, but really setting the benchmark.

This is strictly me talking, I think there have been a few points of frustration. There have been probably too many entryways, too many different kinds of Copilot. All of that now is being really ruthlessly focused in on. Ultimately, I think everyone at Microsoft knows the challenge. It has to be *the* enterprise platform for the AI era, full stop. No question. The competitive pressure is real. The pace of AI is insane.

I'm pretty confident that there's the clarity that is there. There's been a lot of recent internal reorgs to deliver something that is going to be way, way beyond Claude Cowork, probably the best current AI product. I think there are a lot of things that only Microsoft can do. We're already the partner to something like 493 of the Fortune 500; trusted, world class, cybersecurity, clean data lineages in our models.

Just in terms of that trust, that security, that foundation, that embeddedness in the workflows of the world's top enterprise businesses, and the clarity of focus on delivering something that is going to be a meaningful productivity transformation for those companies, I think it's coming together.

TM: We've seen with Zoom versus Teams, you don't have to be first. You just have to be able to get to a competitive product fast enough to then take that market again. And there must still be that incredible pressure because I can't imagine Microsoft enjoys being behind.

MB: No, definitely not. There's a pressure and a need for Microsoft to be self-sufficient in AI, and there's a need for Microsoft to deliver just a huge, huge win for businesses. Can it do it? Yes, I'm pretty sure it can. And, actually, is by far the best placed company to deliver that globally.

TM: And as you said, there's been some fairly well-publicised internal changes recently; Mustafa taking on a new role, head of Superintelligence, and some other reorganisations within the product group. So, maybe just a little bit of illustration as to what those changes mean within Microsoft.

MB: Yes. And, you know, some of it's not announced. We can talk later about superintelligence and what we're trying to build. Microsoft AI is now a fully-fledged frontier lab aiming for the absolute state of the art in AI. And that wasn't always true. It wasn't always in a position to be there. You will have all followed the deals with OpenAI; there's a recent announcement. There was one last year that essentially freed Microsoft to be completely going for the absolute frontier. So, that's been one change.

And then until recently – and this is all a matter of public record – there were a bunch of product

teams working on Copilot. That's now changed, and there's a merging of product teams on Copilot to give a far greater focus and clarity. What I probably can say is, watch for that to now translate through into the product. That's only, maybe, seven weeks old. But over the next six months, you should see marked change that comes from bringing different product teams, owning different products into one organisation.

TM: We'll look out for that. Following the track, and closer to where you started on this exponential compute build, I think the thing that we've probably debated the most at Walter Scott is – and it's evolved our thinking over time – it's this return on investment question around AI. So not just Microsoft but, obviously, you're looking at the cumulative spend plan for this year. The hyperscalers have got over US\$700 billion dollars earmarked for investment. There's more outside of that. Microsoft, US\$190 billion, so very, very sizable amounts and increased that materially recently. So, how would you frame the thinking around this return investment question, either from your personal thinking or Microsoft's thinking?

MB: I'm sure there's probably a different, perhaps more diplomatic corporate line, but here's my take on it all, as indeed a Microsoft shareholder myself. The numbers are eye watering, and go back five years, it's hard to imagine that the hyperscalers would be investing like this. It is huge, but the way I think of it is like this: the AI revolution that is happening now probably has at most three parallels in modern corporate history: the first industrial revolution, the second industrial revolution of electricity, automobiles, and so on, the computer and internet revolution, and now this. And maybe it's akin to the first and second industrial revolutions put together.

And anything less than that level just simply, in my view, doesn't do justice to what is actually happening. We are creating cognitive abundance. We are commoditising intelligence, the thing that has created the entire world around us that is the engine of everything that humanity has done. So, this is an industrial-revolution-level event, potentially more.

If you're a hyperscaler to just say, "Right, we're going to sit this out," is probably a bit of a dereliction of duty. This is about putting in place the foundations of a new industrial revolution that will shape the global economy for the next hundred years or more. That's the basic frame in which I think it's worth thinking about this. The numbers are huge, but the potential consequence is huge as well.

And, just thinking about this on a global level, if all of this investment even just ticks global productivity growth back up to the level that it was in the second half of the 20th century, if we just match the second half of the 20th century over the next 50 years, that unlocks an extra 100 trillion dollars of global GDP

every year. Even relatively modest productivity gains over that period translate into a huge increase in global wealth.

I also think it's worth saying that AI is how we get the next generation of medicines, the next generation of other technologies, the next generation of businesses. If we don't have this AI, probably some kind of stagnation becomes the norm. On this global level of all of the investment, I'm pretty confident that it pays off.

On a company level: Microsoft is the enterprise foundation for almost every big company on earth. If you believe, as I do, that this is this industrial revolution, this is how all of those companies achieve their next 50 years of growth at the same level or massively in excess of the past 50 years. And again, for each one of those companies, that is a huge, huge, increase in what they can do in their revenues. If Microsoft captures a small but reasonable portion of that, it definitely pays off.

That's my view on it. It's a once in a lifetime moment for every technology company. To fall behind now would be a colossal mistake. The numbers are big. It's a wild moment. AI is creating a unique time.

TM: It's certainly an interesting time. I guess, the evidence over the last couple of quarters in terms of the uptick in revenues, largely driven by coding agents – Anthropic, OpenAI, both absolutely flying – and seeing that also in the numbers of Microsoft and other hyperscalers, you can really see that inflection over the last couple of quarters.

MB: Yes. The problem that we face even internally, and this, I think, is replicated across almost every business that relies on code, is that there's simply not enough compute available. Everyone wants to use more tokens, and they're hitting limits. Anthropic are throttling Claude; OpenAI, all of the companies are limiting the amount that they can offer or massively jacking up prices just because there's no compute. Everyone can see that the agents are delivering. They want to use the tokens, but there just simply aren't enough GPUs in the world to serve the amount of demand. There's not enough energy in the world to serve the amount of demand that's out there.

Every time they forecast demand for AI, it's just been blown out of the water. When you see these big numbers, it is based on quarter after quarter after quarter of demand outpacing supply of AI. And this is something that is experienced by people I sit next to who are rinsing through their tokens, building things that would have taken weeks or months in hours, and then getting frustrated because they can't keep going.

TM: I know you've been writing about this, and you talked about the exponentials, the trillion-fold increase in FLOPS [Floating point operations per

second] so far. Would you mind just explaining what you see going forward? Because I know you've been working on the numbers to 2028 or 2030, just how extreme it gets as we go up the exponential.

MB: On my team, we have forecast that by the end of 2028, we'll see another three orders of magnitude of compute globally; potentially up to a 1,000x what we have now. By 2030, it's forecast that we'll be bringing online potentially 200 gigawatts of compute for AI. That is equivalent to the annual electricity of the UK, France, Germany, and Italy brought online just for AI in 2030. So, I think there's one of two things that is true at that point; if that is happening, it is because AI is delivering at a huge scale. We have what I think of as professional or organisational level AGI where the AI is as capable as any white-collar worker or even company.

If that hasn't happened, then I don't think we'll be bringing that on. There's just not a world in which in 2030 without serious gains coming in, we're bringing on 200 gigawatts of compute. In the positive case world, there are still bottlenecks. There's only so many silicon wafers that can be produced. There's only so many ASML, extreme ultraviolet lithography machines that can be produced. So, the 200 figure is the absolute cap on what can theoretically be possible with everything that is already in place now.

But, if you accept that perhaps the scaling laws will keep holding, and that those scaling laws will unlock an AI that is as capable at doing work as basically any human, then we can do it. There will be the compute there to do it.

TM: It's a fascinating setup. We've been speaking to TSMC and ASML in the last couple of weeks, and there is this fascinating supply versus demand tension that is building. We've already seen the prices of memory and many other inputs skyrocket – storage, photonics, et cetera.

MB: Yes. I'm sure you're an audience who are extremely plugged in, but the cost of DRAM memory right now is absolutely skyrocketing. The cost of a laptop or a mobile phone is about to go up a huge amount; not because of anything in Iran, but because it used to be a small amount of the world's memory went to data centres. This year, it's going to be something like 75% of the world's memory going into data centres. If you're somebody out in the world who's buying a \$100 mobile phone, that's probably going to start costing \$300. This is a huge shift where AI demand isn't just a tech story, it's a macroeconomic story that is about to play out over the consumer economy.

TM: It's fascinating, these companies, they're hearing this trend, but they're so used to normal cycles, they're determined not to be a bottleneck. But I think it's inevitable if that Microsoft case plays out, they will be, right; they'll continue to be a bottleneck? It's just

going to be super hard for them to get the conviction to build that capacity, even if it's possible to build that. It's going to be a fascinating couple of years as the exponential hits more and more reality.

MB: TSMC and ASML are not known for overbuilding; they are cautious companies, so what you say is absolutely true. Still, we think that with everything that's there, there is enough capacity should the demand carry on. That is, I think, an important caveat. No hyperscaler is just going to continue to build out at this level if it doesn't see a return that it thinks is worthwhile.

TM: Let's pivot to talk about risk; if you read *The Coming Wave*, there's some incredible sections on different risks. Really, the whole book from my perspective, is this call to action for the world to think deeply about AI, to think about the problems that are inevitably going to arise from AI but also providing some solutions and some hope and some deep thinking around that.

So, let's delve into some of the risks. One area that you talk about is this rise in asymmetry. This ability that AI gives individuals and small groups of people either to do brilliant things, like create products in a very short space of time that are great for humanity, but also nefarious actors – cybersecurity, bioweapons, things like that – these are covered in the book in depth. I'd love to hear your updated thinking, particularly around that risk of asymmetry, and maybe can we focus on cybersecurity given it's so topical around Mythos, et cetera.

MB: It's important to say, I've spent a long time talking about the risks of AI, and the risks of AI go from brain rot on one level to existential for humanity at another level. And they're all real. But I am pro AI, and we can come back to why.

AI is a power booster. It just gives anyone the power of the world's most extraordinary technology. If you can access the model, then what you can do in the world massively increases. And, obviously, that creates a huge amount of positive impact in terms of people can do more great stuff. It also means if you're a bad actor, you can do more stuff because the AI is more capable.

The cyber impact is now becoming extremely real. Claude Mythos, which you may have heard of, is an extremely powerful model that they haven't released publicly, but they have released to a small group of tech companies: Microsoft, Google, Apple, and others. There's public debate, is it real, or is this another hype? I can tell you it's definitely real. Claude Mythos is a cybersecurity event that is extremely significant. And you can see this, because some of the companies that it's been released to – nobody's saying anything – but they have released how many bug fixes they've put out. I can't remember the company, but I saw in the last month

since they've had Claude Mythos, they've shipped 10x the bug fixes of the previous month. So that is Claude Mythos has found ten times the security vulnerabilities in that company, compared to the previous models that they would have used for their red teaming.

Honestly, whatever security you have, upgrade it. Get physical security keys, things like that. Because what's for sure is the progress doesn't stop with Claude Mythos. It keeps going. And so, we're now in an arms race of cybersecurity that is quite unsettling. This is basically a huge risk. What we're going to rely on is responsible rollouts and, credit to Anthropic for rolling Claude Mythos out in the right way, and we're going to rely on AI enabled security.

So, it's a risk; it's something that, genuinely, everyone should take seriously, and those are probably the only two things we can do: responsible rollouts collaboratively and AI enabled security.

TM: When you listen to certain tech leaders, Sundar Pichai at Alphabet, and Marc Andreessen, [they] are basically resigning themselves to a period of cybersecurity chaos until we get to the other end of that where we've fixed all the gaps, fixed all the holes, and actually, the end result being much more secure software platforms. Is that your view of how this progresses for a while?

MB: Again, this is me talking. That seems plausible. Hopefully the good labs, responsible labs, can just stay in front enough to mitigate that. The worry for me is that somebody hacks in [and] steals a model. There was this incident [with] Mythos, there were some unapproved accesses to that; or let's say an open-source model somehow gets more capability than expected, a kind of DeepSeek moment above cybersecurity, and then you have the chaotic world. I think there's a very real risk of it; but we just have to stay focused and, again, this is where a bit of regulation comes in.

If you have a dangerous model, you can't just put it out into the world without some kind of check happening beforehand. The AI [Security] Institute here in this country has been really good about checking things, there is also one security institute in the US. They're doing great work to actually look at models before they're going to be released, understand the risks, and then work with a company or a lab to make sure the rollout's safe.

TM: I think it'd be remiss not to finish with something that you were working on more recently. We talked about earlier, something that's really unique and differentiated about Microsoft and your writing is this idea of humanist superintelligence. Can you just explain that concept because I think it is incredibly logical, but an unusual voice in this community.

MB: What we are trying to build, is what we call humanist superintelligence – so humanist AI. And I think it starts from the perspective of what is superintelligence? What is this thing that these companies have said they want to build? Well, superintelligence is basically something that is pretty much better at cognitive tasks than all of humanity put together, that can go and do anything, that is basically an autonomous being that can go off and achieve all of that.

And, actually, we think, well, why would you want to build that? If you start with humanity and you start with the interests and the future of humanity, it doesn't seem to me like a great idea to just go and build something that is completely general, completely autonomous, that is just better than us on every possible dimension. That just seems like a big safety risk. It doesn't sound like something that is controllable. It doesn't sound like something that if humanity decided to, we could just switch off.

It also doesn't sound like something that is built with human interests in mind. Is this about supporting people, augmenting people, or is it actually about replacing people? And that is the vision of superintelligence that basically has dominated. And some people have been clearer about that, some people have fudged it. But what very few in the tech world have actually said is, hold on, we're going to build technology that is actually all about human flourishing. And as a team, we're putting flesh on the bones of what that means. We're about human flourishing. We're about human control. We're about safeguarding the human future. That means, actually, we don't want a superintelligence that is amazing at everything. Actually, we want a superintelligence that is super intelligent in some domain, let's say medicine, let's say energy, even let's say music; but actually maybe then doesn't have access to the internet, doesn't have autonomy, isn't totally general.

We have this idea that humanist superintelligence stays within the control of humanity forever. If it's not delivering for humanity, it can be changed, turned off, made to do so. It basically says, let's just not go all the way here. Let's try and have a bit more of a sensible idea of AGI and superintelligence. And I feel like that's a common-sense position. I feel like that is probably the position that most of us would get behind. If you go up to somebody on the street and you say, hey, do you think it's a good idea to make this superintelligence better than all humanity, at all tasks, all the time, forever? Most people will tell you, "No, why would you build that?" And, basically, I agree. So, what we're going to try and do is build superintelligence with *most* of the benefits, but far fewer of the risks that we're going to build something crazy.

TM: I don't know if that, to the room, feels like incredible common sense or, why would you challenge that view? But if you actually listen to other

tech leaders, and I won't name names, but it's easy to get interviews with very big players in the industry, which are clearly not on Team Human. They are looking to supersede the human race and even talk about us living on reservations and just wild concepts like this. So, it's good that you guys are riding for Team Human.

MB: This is something that I've been researching recently. When you look at a lot of the motivations for building superintelligence or even just researching AI, you get into some strange places; you have people who want to merge with the AI and upload their minds to the cloud. You have people who are building it because they want to understand the mysteries of the universe. You have people – not that many – but some people who are outright successionists, who believe that they're building the successor species to humanity. And our role is somewhat like the dinosaurs, that we should create these silicon beings and then resign ourselves to our end. That's not what I think is right. I think we should do this only because it improves human lives, human well-being. And if we can't be sure of that, we should reject it.

TM: That's been absolutely fascinating; so thank you so much. Unfortunately, we are out of time. I'm sure you have more questions, so please do try and grab Michael at some point. I think you'll be hanging around for a little bit. Thank you so much for joining us.

Important information

The views expressed in this transcript are those of the participants as of the date recorded and are subject to change. While Walter Scott believes the information to be reliable, no representation or warranty is made as to its accuracy or completeness.

References to Microsoft are provided solely for informational purposes and to illustrate aspects of the research process. Such references should not be construed as investment advice or as a recommendation to buy, sell or hold Microsoft or any other security. Any references to portfolio holdings are subject to change and may not represent current holdings in any Walter Scott strategy.

Certain statements may be forward-looking in nature. Actual outcomes and results may differ materially from those expressed or implied.

1.0_COM003026_26/08/2026

Walter Scott & Partners Limited, One Charlotte Square, Edinburgh EH2 4DR
Tel: +44 (0)131 225 1357 · Fax: +44 (0)131 225 7997 · www.walterscott.com

>BNY | INVESTMENTS

Registered in Scotland 93685. Registered Office as above. Authorised and regulated by the Financial Conduct Authority.
FCA Head Office: 12 Endeavour Square, London E20 1JN · www.fca.org.uk