



SuperDataScience

SDS PODCAST

EPISODE 1016:

IN CASE YOU MISSED

IT IN JULY 2026



Jon Krohn:	00:00	This is episode number 1016, our ICYMI in July episode. Welcome back to the SuperDataScience Podcast. I'm your host, Jon Krohn. This is an ICYMI an episode that highlights the best parts of conversations we had on the show in the past month. My first clip is from episode 1013, where I sit down with Dr. Cathy O'Neill, Harvard Math PhD, former hedge fund quant, and author of the mega bestseller Weapons of Math Destruction, a decade on from that iconic book. Kathy makes the case that what makes an algorithm terrifying is not its complexity, but other features entirely. Listen in. Yeah.
	00:45	We were a lot less vulnerable, I suppose, back then when these automated systems were so much simpler, where random forest or support vector machine is kind of as complex as it gets. I mean, I suppose there's still a lot of room for exploitation there, but now, I mean, even 2016 is when weapons of math destruction came out and a ton has happened since then in terms of AI capability.
	01:12	With large language models, deep fakes, short form video being created by GenAI, surveillance pricing that has largely become dominant since 2016. Do you think we need a weapons of math destruction sequel that is going to talk about all these kinds of high powered AI capabilities that make us even more vulnerable?
Cathy O'Neil:	01:35	I think the things that make algorithms terrifying aren't their complexity or their technological advancement. It's a combination of the secrecy, the unaccountability that nobody in particular is in charge of checking if it's right or in charge of mistakes. And the fact that people absolutely have to use them. They don't have choices. They could be flow charts. They could be linear regressions and often are, are logistic regressions. It don't have to be complicated for them to be terrifying and exploitative. So I mean, a lot of the algorithms that I wrote about in weapons that are still there and there's probably the most terrifying algorithms of all. I would argue there are new



ones. I would say the AI chatbots that tell kids to kill themselves would rank up there as one of the most terrifying of all. So I'm not saying there's nothing new under the sun, but I'm saying the ones that I wrote about are still awful and they haven't been addressed.

02:44 To answer your question though, could there be a sequel? I actually considered writing a sequel. One of the things that changed my mind is that everything is changing so quickly that it would, as I said, takes four years to write a book. It would be completely outdated by the time it went to China to be published and then came over on a ship. It's 12 months later to be sold. And that's why I started my podcast. I literally started the AI Skeptics podcast because I was like, yeah, it takes too long to write a book nowadays. You can't address these issues four years from now. But yeah, I guess I've always considered the complexity to be independent of the terrifyingness of an algorithm, of a system. Absolutely the most terrifying one of all for me was the one called recidivism risk algorithms that sometimes judges were using to sentence people to longer in prison based on risk score of them someday getting rearrested.

03:51 Still being used, still sending people to prison for crimes they have not yet committed and might never commit. And by the way, simple as pie. Sometimes they're just questionnaires and then points are added up. It's stupid how simple they are. But I still think they're unconstitutional. I don't understand how they're used.

04:11 It just doesn't make sense to me. But the power is in the algorithm and the power is what is terrifying.

Jon Krohn: 04:18 That makes a lot of sense. Since 2016, it seems to me like there might be some places where we can now be using algorithms that previously would have been ineffective. So in chapter seven of your book, for example, you mentioned how productivity management models are



optimized for efficiency and profitability, not for justice or the good of the team. And for a long time, things like the movement of hourly workers in warehouses, fast food chains were monitored to enable automation and other efficiency gains. But now with large language models, we can be doing things that we couldn't do before. So for example, Meta announced its model capability initiative, MCI, I think that was this year, just a few months ago, which is a plan to track every keystroke and mouse movement of employees to train its AI models, suggesting that the kind of demeaning, degrading of labor that has long affected blue collar workers on say factory floors or fast food chains is now affecting even the highest paid echelons of white collar work at places like Meta.

05:27 Yeah. There wasn't really a question there, but I feel like you might have a response.

Cathy O'Neil: 05:31 I especially like the way you ended that, which is like I interviewed truckers and teachers and people who had gone to prison or are being denied parole based on recidivism risk algorithms. And they were. Let me just give you a little non-answer to your non-question. Those are the people I was talking to when I was writing the book and I was like, "This is demeaning. This humiliating. This is degrading. It's dehumanizing. Truckers are working for algorithms. They're working for surveillance systems. They're not working for themselves anymore." I used to, by the way, load trucks in high school. And I met a bunch of cowboy truckers back in the day in the '80s. They were just absolutely cowboys. I mean, they were also hopped up on all sorts of different kinds of drugs, but really interesting, very, very interesting, kind men, mostly men and independent, really independent, like independent businessmen.

06:38 And it's very small business, usually just themselves. But that no longer exists, that model of the cowboy trucker. It's just not possible anymore. And that way of life just



being removed and being replaced by something that is much, much less interesting and sexy and human, it was sad. But true. And the teachers that were being fired based on almost a random number generator, similar, teachers have been underappreciated, underpaid because back in the day we had free labor basically from women like Louisa Mary Alcott or Laura Ingalls Wilder. They were just like super smart women. They were only allowed to teach, so they were underpaid. Anyway, the point being that I was working with people who are workers and they had a very strong notion even back in 2014 that when I was interviewing them that like, yeah, these algorithms are working against me. This is something that's happening to me.

07:48 This is happening to me. And then I remember I was asked to give a TED Talk the year after my book came out because it was pretty popular books, as you say. So I went to the TED main stage and I gave a TED Talk and I talked to some audience members and they were so excited about big data. And they were just like, "Oh, this is happening for me. This technology is happening for me. It's going to make me more productive. I'm going to have..." Elon Musk came the same time that I came. And even back then, by the way, I just was like, "This guy is a jerk." And I'm holding myself back because I know you bleep people on this podcast. But people were just super into this idea that Neuralink, they're going to have an automatic connection to the internet and they'll be smarter.

08:42 They're going to be smarter and faster and more productive and things are happening for them. And I just was like, "This is the new divide." And that was whatever, 2017. I was like, "This is a cleavage in our society, which is only going to get bigger." That's what I meant when I said how big data increases inequality and threatens democracy. The cleavage is the people who think the systems are working for them versus the people that



think these systems are working against them or it's happening to them, which is even more sinister way of thinking. But to go back to your point, Jon, and sorry to ramble, but for me, what's happened with AI and the engineers whose keystrokes are being measured, blah, blah, blah, which by the way, that happened to me in the hedge fund. All the keystrokes were being measured then too, just for spying issues.

09:38 But yeah, so white collar workers are getting what blue collar workers have had for many, many decades, which is a humiliation and degradation and dehumanization. And in some sense, I think that's a good thing. Not that it's a good thing, it's not, but it's good to have more people on the side of the workers. It's good to have solidarity to understand, yeah, this is real. This is happening. AI is stealing your ideas, stealing your work. You are no longer necessary because we've already stolen the stuff that we need. Not really true though. That's the good news. Going back to my skepticism about whether digital work is really going to be replaced by AI. No. Some stuff will be, but not all of it.

Jon Krohn: 10:29 Kathy closes on white collar workers now getting the surveillance and degradation that blue collar workers have lived with for decades, though she remains skeptical that digital work will really be replaced. In episode 1007, the founder of 80,000 Hours, Ben Todd, back for his second appearance on the show, takes that premise seriously and asks, "What happens if we do get a fully automated digital worker?" Well, we cover where solid ground is left for human careers, why the bottlenecks then move into the physical world and how fast a robotics build out could really go. If a frontier lab comes up with a fully automated worker that basically can just plop in, can do anything that a remote worker could do over any time horizon. It can do years of work kind of independently. It checks in at sensible intervals. It does everything. It's imagining that somebody that you've only



ever had Slack email and Zoom conversations with, there's no technical reason why you couldn't have that be something completely automated in the future.

11:36 In that scenario, when I asked you about where there's going to be solid ground for years to come and you were like, "Well, just always focus on that little bit. You're going to be able to drive a lot more value." But in that scenario where there's this fully digital worker, I mean, it makes the solid ground feel pretty narrow, doesn't it?

Ben Todd: 11:56 Yeah. I mean, just quickly, there could still be legal issues. An AI wouldn't be able to own a company, for example. And there could be liability issues like that. So that would still. But yeah, I mean, if you actually get the full digital remote worker, then yeah, the bottlenecks move to these things that either have to be done by humans for some reason, such as legal ownership or maybe just the consumers have a really strong preference for it to be done by human. But then there would still be the physical bottlenecks, right? So there's still jobs that require physical presence.

Jon Krohn: 12:34 Somebody's

Ben Todd: 12:35 Got

Jon Krohn: 12:35 To put those GPU racks together for now.

Ben Todd: 12:39 Well, yeah. I mean, this is kind of what's happening is a lot of construction and energy and data center building. These are big growth areas and they're also complimentary with AI.

Jon Krohn: 12:51 For sure. But I've also got to believe that if we have AI systems clever enough to be complete digital workers, then it's not going to take those digital workers very long to be figuring out how to be automating the hardware stuff too, the physical installs. Again, I suppose-building robots. Yeah, exactly. Yeah, that's what I mean. Having



physical embodiments. There's still some testing that needs to happen with physical embodiments in a way that software scales more easily because you can just copy the model weights and you can do that almost for free and almost instantly. Whereas with robotics, it scales a little bit slower because you have to manufacture something, you need to test it. But there's all kinds of ways that you can use simulations in order to be able to train robot arm model weights much more rapidly than from physical real world data alone.

13:49 You probably still want to do some final testing in the actual real world before rolling out some product, but some hardware product. But yeah, there's a lot of. Yeah. Even the construction of these data centers and stuff you could imagine being done by robots in a world, not maybe just a few years after we have these fully automated digital workers. You could imagine them starting to make a really big impact in the physical world as well, right?

Ben Todd: 14:18 Totally. Yeah, I think people sometimes are a bit too. They think this might. They kind of assume this will be a very long process, which it might be, but I have a Substack post about how quickly could you scale up robotics. Because imagine in. You also need to remember in this world, if you actually have a digital remote worker, then the physical bottlenecks become just the whole bottleneck on the whole economy. So you then have this potentially massive mobilization to try to solve that bottleneck by the world's biggest companies. And so what they might be able to achieve in that type of scenario could be pretty dramatic. And one analogy you could look at is something like how quickly was airplane production ramped up during World War II? And that was partly done by converting car factories into plane factories. And so one very rough estimate is if you converted all the car factories into robot factories, how many robots could you make?



- 15:16 And just on a kind of mass per mass basis, it would be something like a billion a year. And in World War II, they were converted in a matter of years.
- 15:26 Obviously, robots are significantly more complex. The hands are much more complex than cars. So maybe that's a false equivalence. We do have a lot of industrial capacity that could be converted and it could go pretty fast, I think. I mean, also this is a world where you have AI aiding you in all of the steps. We have now a full. Well, I think it would effectively be superhuman. This is another thing people don't understand is once you get a human level digital remote worker, you're basically getting superhuman abilities immediately after because you can speed them up 50 times. A day for us is a month or two for them. And there's all the other AI advantages. They can share their state space across all their copies. So any learning can be just immediately propagated to. Yeah.
- Jon Krohn: 16:21 Yeah. In the same way that autonomous vehicles already have lower crash incidences than humans, but humans aren't constantly learning how to drive more safely. Whereas based on more data being collected, based on improvements in algorithms and sensors, autonomous vehicles are always improving. And that 50X thing, I think the point you're making there is that a task that could take a human a couple of months could take hours or days for machines because yeah, you could have 50 agents working in parallel on different parts of a problem.
- Ben Todd: 16:58 Yeah. Another big aspect is the kind of coordination aspects where companies are pretty inefficient because it's like there's a lot of communication overheads between all the different people. But if you have AI workers, they can all kind of communicate instantly across the firm. And you can have the CEO effectively personally supervise every single worker, which is obviously a huge bottleneck now. So these AI firms could actually, even if they just had human levels of intelligence, just through



being able to coordinate much better, they might be able to move way faster than human firms.

- Jon Krohn: 17:41 Those are forecasts at the scale of the whole economy. My next clip brings things right back down to individual people. In episode 1009, venture capitalist and five-time entrepreneur Steve Mock walks me through the patterns emerging from the stories he collects at his website, aisavedme.org. Chief among them that the people getting the most out of AI in healthcare are not asking it for advice, but using it to become far better informed advocates for themselves. Stick around for the story of the Australian dog whose life was saved by AI. And you've mentioned to me that one of the big patterns that you've seen emerge from the submissions on aisaveme.org is AI being used for healthcare. Do you want to dig into that?
- Steve Mock: 18:21 Sure. So what I'm starting to see now, because we're getting lots of stories on the site, is we're starting to see now is patterns within categories. Let's go back to my 84-year-old dad. I'm trying to answer his question, how do you use AI? So there's spot use cases. Oh, you can use like this, you can use it like that, you can use it like this. And then I'm starting to see what I'm going to call a best practices layer appearing in different categories. So in healthcare, for example, I've got a number of stories and nobody is saying, "I went to AI and asked it for healthcare advice and did it." By the way, there's people that do that. We read about it in the press. And I think this is where the education piece comes out, which is, and this is a very simple example, but what all the people in my stories are doing is they're going to AI and they're using it to educate themselves on the subject matter of whatever their predicament is.
- 19:22 And they're learning how to then ask thoughtful questions to their medical provider.



19:28 And they're not getting better at healthcare from AI. They're getting better at asking educated questions to their healthcare provider. So again, it's helping us become our own advocates, which we have to do. And that was otherwise a very expensive proposition before the LLMs existed. So the don't ask AI for healthcare advice. We get it because we're in the thick of it, but a lot of people don't get it. We're reading stories in the paper where bad things have happened because of that. But the best practice, which I'm hoping rises out of the stories in the site, which I'm hoping we get the word out more, it's these types of patterns, which is use it to become a better advocate. Don't use it for healthcare advice. Makes

Jon Krohn: 20:10 A huge amount of sense. I already relayed a story to you socially a month ago when you and I met up around a relative of mine who had been diagnosed years ago with Parkinson's disease. And his primary caregiver, whom I'm hoping will actually submit to aisame.org. But until she does, I don't feel like I should mention people by name. But

20:37 The point that you're making ties into how nobody is more invested in your health than you and your caregivers. That's correct. The healthcare providers, the vast majority of them are great, but they have so many conflicting things on their attention. And they're trying to get as quickly as they can through your case and onto the next one and get through their day. And so they're not going to spend as much time as you can as the patient yourself or as the caregiver for a patient because if something is really big, a life-changing thing like a Parkinson's diagnosis, it's kind of every day. There's things that you need to be dealing with. And you can start to notice things that are different from maybe what your healthcare provider has said or the guidance that they've given. And so yeah, this relative of mine, it turns out that after years of Parkinson's diagnosis and being



given Parkinson's medication, through increases in medication actually making symptoms worse,

21:41 That caused a conversation to happen with AI and say, "I'm seeing these strange. Should this be happening? It seems like things are getting worse, not better." And that led to a situation where they were able to get a completely different diagnosis. They had to go to a different physician. And now they're still working through exactly what is going on, but it seems like it isn't Parkinson's. And so the treatment is very different. And so it seems like a path that looked quite scary and seemed like it was getting worse quickly. Actually, it isn't the path that you need to be going down. That wasn't the medical problem. It was a treatment problem.

Steve Mock: 22:18 Yeah, you're reminding me, and that's what we're also seeing. One of the stories I have on my site is a woman whose dog was chronically ill. And she'd gone to a number of doctors and the doctors didn't have an answer to what the problem with the dog was. She took all the symptoms, put that into an LLM. And then the LLM came back and said it could be this, this disease that she had never heard of. And then she took that back to her doctors and they said, "Oh, we can test for that." And then that was what the problem was. So again, it's helping be your own advocate. It's a resource to help understand better and work with your practitioners for that. Don't

Jon Krohn: 23:00 You also have a wild story on your site of somebody maybe in Australia using AI to come up with a treatment for their dog? So actually inventing a new treatment with the help of AI.

Steve Mock: 23:12 So this is Paul Cunningham, and he's a entrepreneur data scientist in Australia. And I had the pleasure to connect with him over this story. And this is well written about. I highly recommend people go read this story and



there's been a lot written about it. His dog had cancer, had tumors, and the dog was resisting the treatment that they gave. And he then got biopsies of. This is a pretty crazy story. This is not what your average person is going to do. He's an exceptional. He's a dog lover. And he said, "I'm going to do everything possible." He got biopsies of the tumor. He then went and had their DNA sequenced. He then used those sequences to go and find what possible treatments out there could match. He found an experimental treatment that they could actually match that particular cancer. Then he went and tried to buy it.

24:14 And because it was experimental, they refused to sell it to him. So he then found out that you can actually make your own treatments through RMNA vaccine labs. And he went to one of those, had his own custom treatment made. Then they applied it to the dog and within a month, the tumors had shrunk away. Wow.

Jon Krohn: 24:31 And

Steve Mock: 24:31 It's an incredible story. And what's even more incredible is all but one of the tumors shrunk away. And it turned out that it was a different type of cancer. So it wasn't reacting to the other type of cancer. So he did the whole process again. And then last I spoke to him, he had just given that second treatment to his dog. But he talks about now, he said last I spoke to him his dog is thriving. And it's an incredible story. We're

Jon Krohn: 24:59 Rounding up a great month with episode 1011 in which Dr. Katherine Williams, chief data officer at the nonprofit Candid, takes on a question I get asked all the time. "Does a deep understanding of the underlying mathematics of machine learning still matter? Now that large language models are getting so good at exactly the math and programming our field used to prize. "Catherine's answer walks up and down the hierarchy of complexity from arithmetic to systems thinking and lands



on what she believes will separate the best data professionals from everyone else. Having a really deep understanding of the underlying mathematics is still really useful today. I'm guessing the answer is yes, because you talked about, for example, being able to understand ML papers in detail by being able to work through them. And obviously if somebody doesn't have a rigorous math background, they can't dig through all that.

25:48 But we're also in this interesting time now where large language models are getting particularly good at. If you look at the meter charts of capability, those are based around capability on programming tasks, math tasks. It's a very interesting time that we're in right now where I wonder how those technical skills, any of the ones I just mentioned. I said engineering seems to have supplanted kind of mathematical background in importance in our field, but both of those things, programming and math are the two things most vulnerable to disruption by large language models. So anyway, very long question. My apologies for that.

Catherine W.: 26:35 Yeah. I think about this a fair amount in different flavors and I don't have a great answer for where it's all going. I mean, I think it's clear that there's a hierarchy of different levels of complexity and these things build on each other. So there's arithmetic that gives rise to algebra, then you have linear algebra, and then you have systems of things that lead to machine. I mean, these are bad examples, but there's just layers and layers of conceptual hierarchy involved conceptually. And then also on the physical side, right? You have electron, you have circuits and you have chips and then you have hardware and then you have motherboard. I mean, you can walk the stack of complexity. And I think studying any one particular piece of that can be very valuable because each piece in that chain plays a role and has gotchas and has value to being deeply understood and researched and so forth.



27:32 And as machine capabilities start being able to do most of the work in those areas, it's less necessary to go deep. And so you then move to the next higher level of the hierarchy and focus your attention there or on assembling the pieces. But that doesn't mean that understanding lower levels isn't still really valuable. And in fact, it's really important. In my mind, it's parallel to the argument of should you teach kids arithmetic when they all have calculators? Well, yeah, because you need to have that sort of fluency with conceptually, even if you're not going to do the arithmetic yourself, you need to have the conceptual fluency in it. Sometimes I think about it as like we're training our own neural networks so that we have the right subsystems to then create the abstractions, to create the abstractions on top of that, that then lead to the right understanding of the world.

28:20 So anyway, the transition from sort of math to engineering to me reflects one movement sort of up the hierarchy. And now that we have LLMs, you can do the engineering part, what's the next level on top of that? Well, presumably some kind of a systems view, but you're still going to need to be able to dive down into the different layers and understand how they fit together, I think. So I think if anything, the best data professionals going forward are going to be ones who can really move up and down that chain and build their mental models and continue to update them as they learn over time rather than specializing in one particular slice. All

Jon Krohn: 28:58 Right, that's it for today's ICYMI an episode. To be sure not to miss any of our exciting upcoming episodes, subscribe to this podcast if you haven't already. But most importantly, I hope you'll just keep on listening. Until next time, keep on rocking it out there. And I'm looking forward to enjoying another round of the SuperDataScience podcast with you very soon.