

SDS PODCAST EPISODE 940: IN CASE YOU MISSED IT IN OCTOBER 2025



Jon Krohn: 00:00 This is episode number 940, In Case You Missed It in October episode. Welcome back to the SuperDataScience Podcast. I'm your host, Jon Krohn. This is an in case you missed it, episode that highlights the best parts of conversations we had on the show over the past month. To start things off, in episode number 933, I ask Sheamus McGovern, founder of ODSC, the world's largest data science conference, if we need to radically rewire our professional skills to keep step with the rapid developments in ai. You mentioned earlier in this episode about how you've been doing this meetup tour across the United States. You mentioned Chicago, New York, and before we started recording, you mentioned to me that when you speak to people, a lot of these people you speak to are concerned about the pace of change due to AI advancements. And so you talked to me about this idea of rewiring professional skills. So what does this mean? How can people be coping with the pace of change that is coming about because of these rapid AI advancements and proliferations?

Sheamus McGover...: 01:08 Yeah, I got so many of those questions and it's almost like, and I've had to do this myself. You have to think about it in two parts. One is you have to get comfortable with the speed of change. And I think most of my career, I've been comfortable with that. That's one of the reasons I left very large companies and started my own startup. I want to do things quicker, faster, so it's not for everybody, but you do need to get comfortable with that speed of change. And there is actually an advantage there as well because it's been proven throughout time, and you've seen this over decades and centuries, that the pace of change always outruns most people's ability to absorb it. And that of course leads to a lot of angst, but also leads to opportunity. So if you are one of the people who learn how to deal with that, and what that means is you have to be comfortable with this continuous need to rewire.

02:17 And I think that's what people have to understand, that the speed itself is the new challenge. And this is something I've been kind of studying a lot. There was a study, I might send it to you for the show notes, but I did read this study in AI exposed industries, and I'm not just talking about data science now, but AI exposed industries are jobs, rather, the skills turn over something like 30 something percent. So more or less, I dunno if it's compounded or not, but more or less every three years, I don't quite believe it, but even at low case, the skills requirements turn over and you can kind of see that people now have learn prompt, engineering, vibe, coding, and that pace of change, it leaves a lot of gaps. Your company may be moving fast, you're moving slower, you are moving fast, the company's moving slower, the industry and stuff like that. And then I've listened to a lot of your podcasts, Jon, of course, and you can see as well that compute is doubling every six months. That's having a big problem. And so getting comfortable with the pace of change is important because I started listening to this show last year about the history of the universe. It's my AI detox program, great show history of the universe and

Jon Krohn: 03:41 What's it called, the history of the universe is the name of the

Sheamus McGover...: 03:43 History of universe is on YouTube history of the universe. It's my AI detox, it's about astrophysics and astronomy and all that kind of stuff. But believe it or not, until about four or five years ago, I didn't know the universe was expanding. So I always think of AI skills like the universe. They never stop expanding. And just like the universe, we don't know what it's expanding into, what the universe is a bubble or there's multi bubbles and all that kind of stuff. But I'm going off on a tangent here. But anyway.



- Jon Krohn: 04:08 No, that was a good analogy. I like that. I hope we turn into, I hope we turn what you just said into one of our animated shorts for this episode because that's a great visual. This idea of just like the universe expands skills are expanding the quantity of them like we talked about earlier, data science branching off into AI engineer, and lots of other more specific paths. And now AI engineer will branch itself into lots of other sub careers. So yeah, I think you've nailed it with that analogy. Anyway, I've taken you off track now.
- Sheamus McGover...: 04:45 Yeah, yeah. And look, I've been in the industry three decades in tech, FinTech, whatever, skills never go away. There's just more of them. And that's why it's just amazing when people get so concerned, well, this is no longer being needed. They still need COBOL programmers. But anyway, back to the rewiring. Yeah, so when I talk to people about rewind, they're like, okay, well what does that mean in practice? And yeah, it's such an important thing. I kind of take an obvious, very optimistic, almost, I would say hardcore view on that. I really think, and I'm doing this myself for myself, I think AI is moving from, we were moving away because back in 2015 we were using data science, machine learning, and even AI is a tool. It's moving away from AI being a tool to ai, AI as being a collaborative partner. And I do think those collaboration skills you build now will help you over the next decade because that helps deal with a, can I build something now that can future proof? So I say you want to wait or you want to build.
- 06:05 That's kind of important because when we look at ai, it's quite clear now it's either going to be AI is either going to automate or automate and both of those present opportunity. So I really think that when I talk to people about rewind their skills, first and foremost, stop worrying about your skills being replaced and start thinking big picture. Now, what's possible with ai? I



remember you said this in one of your podcasts. You said something like With AI, you can do work. Now that was previously impossible, and that's absolutely true. I'm doing stuff that I would never have done without AI before and forget about the role of a data science or a machine learning engineer for a second. I talked to a lot of startups in my other role, and also I saw it at ODSC. All of a sudden we had these people from sales and marketing showing up and they're all talking about agents, you want to learn about agents.

07:00 If you think back, maybe a startup yourself, when sales was a decade ago, right? In sales there was you basically had an account manager, an account executive. Now in sales you have a lead gen specialist, you have an SDR, you have an account executive, an account manager, different roles by the way, you have a sales engineer, customer success person, a revenue officer. If you take ai, a lot of those roles can be rolled up. And AI can either augment or automate those and then think about that salesperson. And let's say you were just doing SDR or you were just doing account management or you're just doing sales engineering, you can now do a whole lot more, but you're going to have to rewire your skills because of ai. And again, I've been studying this a lot, and the more questions I get about it, the more it's kind of a continuous loop, the more I kind of study it.

07:57 And as you know, we have our own podcast, which we need to have you on our back on. And we had, was it Robert Brennan? Sorry if I'm mispronouncing his name, but he's the CEO of Open Hands, which was an open source version of Open Devon, which allows you to automate your work. And I asked him the question, shouldn't people be worried about this replacing their jobs? He's like, look, Seamus, most work today is drudge work. And I really started to research that he's right. If you think about the average person in office, they're

looking at emails, they're doing admin, it's mostly drudge work. And there's this whole productivity paradox and automation product paradox. Even though we've got productivity with automation, the problem with the automation paradox is automation still needs oversight, it still needs judgment. So yeah, I think the new roles are going to be, as I said before, you rewire your skills. Going back to data science and ai, less about building models from scratch, more about designing workflows, managing supervision, evaluation and less to be builders and more orchestrators are less to be building from scratch and more orchestration,

Jon Krohn: 09:24 Orchestrators in an expanding AI universe sounds like a great way to rethink our approach to white collar work. And indeed, many companies are actively encouraging their employees to work alongside AI applications. In episode number, Gurobi mathematical optimization guru, Jerry Yurchisin tells me how Toyota optimize their planning process to manufacture their vehicles. Something else that you have for us, I think that's completely new since your previous appearances on the show are some interesting new real life use cases of mathematical optimization. So you have of course alluded to some of them. We've talked about the burrito optimization game or as a toy example or the new guro bean coffee example. You've mentioned that application areas like supply chain logistics, those tend to be areas that use mathematical optimization a fair bit. But I'd love to dig into a few more cool real life use cases that have cropped up in recent months.

Jerry Yurchisin: 10:26 So we recently had what we call the Gurobi Decision Intelligence Summit. It's our fancy sort of event that we put on ourselves. We invite customers, we invite prospects, we invite anyone who's interested in learning more about optimization. We invite you to come. Last year, it just finished up a couple weeks ago, we were in

Vegas, and so we're bringing in super cool customers that are doing really cool things. And we had a couple talks that I thought were,

- Jon Krohn: 11:03 I like how you had to pause there because you're like, company names go into your head and then you can't say them. So we get some pause, really cool companies.
- Jerry Yurchisin: 11:15 But I will mention a couple, there's a couple that I can mention. There's some that I can't. Sadly. Again, we're the best cut secret in decision making. I guess that's what, if you're going to come away with anything, optimization and guro is the best kept secret because people don't like to talk about us because yeah, why would you spill the beans? But there are two presentations that I really liked. One was from Toyota, and they're talking about how they used optimization for planning of vehicle manufacturing. So they're getting sort of demand forecasts of like, okay, this is the number of this type of vehicle that I expect that customers would want in this region at this time. So you can sort of see if you're thinking about by region over a certain amount of time, the whole sort of fleet of Toyota vehicles that they offer, that's a pretty big problem.
- 12:18 And now you're thinking about, okay, manufacturing that, how can I best manufacture these things, these cars at minimal costs and everything. You sort of see all of the small things that trickle into making a car. It's a very complex process. So they ran through how they're building tools and there is an aspect of LLMs and natural language in this as well. But they allowed their planners to interact with an optimization model that an optimization team built this optimization model, but they allowed their planners to interact with that and do scenario tests and what if analysis on all of these sort of things. I'm like, well, what if the tariffs on this particular thing, what if tariffs go up by from 0% to 10% and then next week they're 80% and then the week after that

they're back down to 10, and then sometimes they're 30%.

- 13:28 It's an insane time to try and plan long-term manufacturing right now, it's insane with all of this sort of fluctuation of particularly tariffs, but they had a tool that had optimization in the back, had sort of an LLM interface where the planners can really interact with this and say, okay, well what if tariffs are this? Or what if my supply of this thing was cut in half or something? It's interacting with the optimization model in a very natural way and getting all of these sort of cool scenarios and really being able to understand, okay, what if this happens? What should I be doing? How should I be manufacturing things at some macro level and really making decisions that will impact the company. It's just providing a whole new way to access optimization to people who don't, they're not going to be writing the models, they're not going to be doing any of the Python coding, but these are the people who are making the decisions, who have all that have all this sort of SME expertise, all this business expertise, all this foundational knowledge of I actually know how to plan manufacturing for cars and stuff like that.
- 14:47 I know all of this, I don't know, optimization, but now this group at Toyota, they did an exceptional job of blending the two and letting people interact with that. So that was one super cool case. And the other one is with total wine. The other ones I could mention is total wine. And it's again, a similar problem of how can I, it's a similarish problem because kind of supply, but it's essentially, if you think about what a total wine store is, it's a massive store that has all the beer and wine that you could ever want, anything you're interested in finding. And depending on state laws, there may be liquors and stuff like that too.

- Jon Krohn: 15:33 And here I thought it was a platform for getting complete complaints.
- Jerry Yurchisin: 15:38 I love it. But what I really liked about their story is if you think about the complexity of decision making that can happen within something like that, it's like, okay, well, I buy a bunch of beer, I buy a bunch of wine. But you're sort of thinking again about complexities and in the presentation, the presenter is talking about, okay, I want to buy just one brand of beer or something. The choices that you have in just that single sort of brand is pretty massive. Am I buying massive cases? Am I buying individual six packs? How am I buying cases of 24 cases of 18? All that sort of stuff. When am I getting them? How often are they coming? How often are they arriving and everything like that. And then now you think about that for pretty much every beer that exists, particularly in North America, or are you importing them every wine?
- 16:41 It's a massive, massive problem and not easy to solve. But what I really liked about this problem is the Toyota folks that I just mentioned, and a lot of our customers, they have what we call operations research expertise in-house. Even the Toyota example, the person who presented it did not have the traditional background of our common customer. He is an AI person, but had some mathematical chops to him. And so it was not super, he took to it a little bit faster than I think some would. But total wine folks, they were a team of data scientists. They were people who did not have a traditional sort of operations research background, industrial engineering. Those are some of the common degree types that people have who have been exposed to linear programming, mixed synergy programming, the mathematical optimization things. That's where you typically learn that. But these were people who were, I'm a data scientist, been doing that for a decade now.

- 17:50 Oh, we have this new problem type that we're trying to solve, machine learning's not cutting it. What else can we do? Oh, okay, I've learned of mathematical optimization. Now we need to actually do it. And so it was a total success story of taking a team of people who did not really know how to do this right away, understanding their learning, their pain points and stuff like that, understanding what worked for them and what, it was just a great story to hear that this stuff, if you're listening to this now and you're like, oh, well, I don't have time to listen. I don't have time to learn all of this, or I don't know the benefits of should I just hire or something like that, but that's also complicated. It could be time consuming and blah, blah, blah. It can be done. You can build a team that can take care of this, that can do this at the scale. And I think this is where a company, this is why I love working for the company I work for, is we don't just hand you the software and say, good luck, have fun, as long as your check clears, blah, blah, blah.
- 19:03 We're not going to talk with you. We have an exceptional sort of support team that helps you with this. So if you get stuck, not stuck with like, Hey, I don't know how to build my model stuck, but like, Hey, this is taking a lot longer than I thought to run. Or we're getting these error messages, or we have issues with this or that. You have people when you submit a ticket with us, you have someone with a PhD in optimization or decades of experience that looks at that and thinks, here's how I can help you. So they leveraged that and they used our support system to really help them, and now they're saving, I don't want to mischaracterize the number, but it's a lot of money and they're being able to reinvest it then. And that's really great about these projects, these optimization projects is, yeah, you're saving money typically, but it gives you an opportunity to reinvest and make things better elsewhere. So those are a couple of

really cool customer stories that I was able to hear. And there's tons more though. Tons, tons more

Jon Krohn: 20:23 Saving money is of course one of AI's foremost benefits to corporations that want to improve their margins. But anyone who uses AI must also stay aware of how they are using systems, these systems and tools as technologists. It's so important to keep asking ourselves, how can we use AI to build a better, fairer, more equitable world? On the podcast, we frequently cover how AI comes with ethical risks that have to be weighed against its promising productivity and efficiency gains. In episode number 935, researcher broadcaster and author of the bestselling book *Technology is Not Neutral*, Dr. Stephanie Hare joined us to discuss her thoughts on how we can install ethical boundaries for our AI use. On the note of developing your book and coming up with these ideas of how technology ethics are treated not just in the West but all around the world, something that you've brought up a number of times is the idea of whether we should have something like the Hippocratic Oath that they have in medicine for technology. And so it doesn't seem like that's, I don't know, it doesn't seem like it's probably a practical thing that we're going to have an international technology Hippocratic Oath come about. It's a nice idea, but so maybe instead of a symbolic oath, are there practical, non-negotiable checkpoints that maybe should be embedded into tech product development, life cycles or some kind of tool set like a Swiss army knife that technologists could work with that maybe is enforced in some way and isn't considered to be a luxury?

Stephanie Hare: 22:06 I think that you've hit on the rub of it, which is the enforcement question. The reason I liked the Hippocratic Oath, by the way, is not because it's like a mandatory thing. Not even all medical schools around the world required that now, and it hasn't always been required for doctors. And it was actually recreated or rebooted, if you

will, after the second World War because of course, as we all know, the Nuremberg trials after the Second World War, there was a special doctor's trial because doctors were actually very instrumental in the Nazi regimes murder of many citizens of several European countries. And they had a special trial for that. And so that led to a sort of reckoning and a crisis within the medical community after the war, which was like, how is it that a bunch of people who are supposedly trained to help keep people alive and indeed healthy and thriving, how on earth were they among the first instruments of murder in a tyrannical regime?

23:04 And I was really fascinated by that. My second area of study was history and specifically World War II history. So I was like Jesus. And they revisited the training of doctors because of what happened in World War ii. That reboot came as a response to an acknowledged universally discussed around the world problem of horror. And I was fascinated by that of the way that we think about trust. Doctors tend to be quite trusted, put a stethoscope and a white coat on them and you're like, oh, you'll do what they say. It's very difficult for a lot of people to push back against a doctor. They have more than us, et cetera. And often when you approach a doctor, you're unwell, you're injured, you're sick, or your family member is. So you need to know you can trust them. So I was thinking about those sorts of concepts, the historical reality of trusted, intelligent people betraying that trust in the worst possible way that they possibly could. How do you then come back from that? How do you restore trust to a profession? Why do some medical schools do something like a Hippocratic Oath and some don't? The fact, by the way, that the original Hippocratic Oath versus what said today is largely rewritten. So what

Jon Krohn: 24:24 Was they don't do it in Greek.

Stephanie Hare: 24:26

No, a lot of them have rewritten it, and I kind of like that. It's basically just the first one is first do no harm, which I think is totally appropriate for technologists to embrace as well. And then second, which is the mission statement in my book, is like, how do I maximize the benefits and minimize the harms, which I personally think is a bit more realistic for utilitarian way of thinking about it, which is there's going to be some harm. Maybe you cannot make the omelet without breaking some eggs. So fine, choose it, choose it mindfully, build it in, have a discussion. It could be democratic. We should all be thinking about this. That implies that people have to be around the table. There's knowledge, there's consent, blah, blah, blah, all that stuff. So that was the only reason I was thinking about it. And the reason I liked it for the medical establishment and thought it might be useful for technologists is precisely because it isn't enforceable.

25:19

It's not about getting a driver's license. You're not allowed to drive your car unless you have a driver's license and insurance. And if you don't have those things, you could get arrested, sued, et cetera. This is more like this is part of joining this community. It's an ethos and it's a sign, I would hope, in the best engineering schools, the best business schools, et cetera, that we teach ethics. And indeed that is actually true in lots of professions. So lawyers have this, accountants have this, civil servants have it here in the UK. The civil service ethics code is really serious. I have several friends who are several servants here, and I really admire them. Their sense of commitment, something larger than themselves is part of their professional training. So I think it would be lovely. This is just my own take on it for technologists to have that in their formation and for them to think about it a lot. If we treated our careers as a vocation, why do you get out of bed in the morning? What are you building?

26:25 That would be something that I think could help, not just with how we design and live and create, but also for our relationship with everybody else, the users of our products, our customers, but who are also our family, our friends, et cetera. So it's just an articulation of the value statement, but I don't think we need to add more regulation to it in the sense of you can't code unless you've done this thing or you can't create something unless you've got, the world does not need that. You don't have to be regulated to do the right thing. You could just decide to not be an earth. Yeah,

Jon Krohn: 27:02 It's kind of this idea, even when you said the first line, I guess, of a typical Hippocratic oath of the first do no harm. It's interesting how with technology, often the primary incentive is first make a profit. It's like our first generate a RR.

Stephanie Hare: 27:22 Well, is it though? I would say that's for companies, that's for a lot of people. Sure. But a lot of people are not just tinkering or necessity is the mother of all invention, the person who invented the washing machine or what. I'm just looking around now. I'm like everything in my house, suddenly

Jon Krohn: 27:42 Toilet

Stephanie Hare: 27:42 GO tool. Yeah, you're usually doing it to solve a problem where you're like, God damn, I cannot take this anymore. I want scissors for left-handed people. Instead, I know the world is mainly right-handed, but there's a whole crew of people who are not being served and they can't scissor things without hurting their hands. I shouldn't invent it. I think it's often hopefully coming from that. Yes, there are people who always start with the profit motive first, good for them. But I think a lot of innovators are more, they're problem solvers and then they're like, oh man, if I did this, I can make bank. Why not? There's nothing wrong

with that. But I think the best stuff comes from solving problems.

Jon Krohn: 28:21 From transpositions of the Hippocratic Oath, we moved to multilingual models with Dr. Adrian Kosowski in episode number 929. Adrian explained a new way AI capabilities could be simply concatenated together in an LLM. Something that I found fascinating about your paper, about your BDH paper is you were able to concatenate literally just like a concatenate operation. You could have one neural network trained on one language, let's say English, and you could have another language trend on let's say French in honor of the mackinaw here and with your architecture. And this seems like a rare thing to be able to do with an architecture that could be the building block of a large language model. You can just concatenate those English and French language models together and because of the sparse activation, it just works and it's a multilingual model.

Adrian Kosowski: 29:21 That's the spirit. And I think it touches on so many different aspects, which I think are good to highlight because it's something, it's something new. It's new in many senses. As I mentioned before, the transformer while obviously being an amazing breakthrough in the focus of machine learning and AI in general does have its limitations in the way we understand its scaling. So if you have two transformers and you put them side by side, there's no really clear way how to connect them in BDH versus much easier in the sense that the model scales in one dimension, we call it the number of new ones N , and it's like a size of a Bain. And then if you want to put two such bains together, you can do it depending on what you do, it'll be a little bit like a mix of the skills that you had, or you can also do some post staining for the combined Bain and make sure it coordinates properly. But definitely if you just put for Bain side by side, you have a model which out of the box has understanding for different

languages or is able to map them into concepts in English, for example, and to work with them.

- Jon Krohn: 30:34 That is very cool. Alright, so with all of these incredible novel capabilities of BDH relative to transformer, so the positive sparse activation that we've talked about, this ability to concatenate that comes out of that, the energy efficiency that comes out of it and compute efficiency that comes out of it. Where are you today? It kind of sounds like you've, with this paper with BDH, with the baby dragon Hatchling paper, we're talking about a billion parameter model, which is about the size of GPT two from OpenAI, which is now some years old, and it performs comparably to GT two despite requiring far less compute.
- Adrian Kosowski: 31:18 Just to reassure readers, listeners, to this point, we are looking at models which at a given scale are on par with models of a given scale. So really it's given all the focus but has happened in the state of the art. We use that progress obviously as the one B models that we produce are comparable or outperform the one B models out there. The kind of focus, and the reason why we focus on this one B scale rescale for demonstrations is that this is a scale at which we are able to achieve instruction following and to start testing other capabilities of a model which is able to actually follow instructions and to have a basic capabilities that we would expect a language model. And this is really for the ease and speed of experimentation. There's nothing particularly stopping us from releasing a super large model like in the 70, 80 billion scale larger. The kind of question which is super pertinent is why do it? Because if you are in the world of language models, just language models versus a certain market, which we could call a bit of a commodity market for the kind of chatbot like applications,

- Jon Krohn: 32:48 Discussions and so on. So your clawed, your Gemini, your chat gt, they're all, they're competing in the same space.
- Adrian Kosowski: 32:57 I think a switch that most of us are most aware of is if you are working with a reasoning model or not, usually you are kind of explicitly aware of the switch, especially with models like GPT versus a 103 with Claude, et cetera. You have this option to go into reasoning mode. And this is the place where we don't want to just yet launch a non reasoning model, which is super large because there's actually, it's not our objective here. What we are doing is we are entering reasoning models, we are entering it from the moderate scale obviously, but this is a scale where we can display the advantage of this architecture. I see. See, notably, yeah.
- Jon Krohn: 33:49 So yeah, so the most promising avenue for you for moving forward with this baby dragon family is into reasoning models. So models where you don't just have tokens output being spit out to your screen immediately, but there's multiple phases of reasoning happening in the background, refining your answer, ensuring accuracy. Yeah, that's where you see the most potential.
- Adrian Kosowski: 34:13 That's it. Lots of consideration, lots of interspection. And also something that we see as extremely pertinent is the ability of reasoning models to work with contextualized inputs and to process them. So if you think of baking for barriers, the limits of 1 million token context, but you have reasoning model which goes through billions of tokens of context. Here you're in a space in which you can, for example, ingest a contextualized dataset private to enterprise, like a documentation of an entire technology, which is like 1 million pages of paper, 1 million sheets of paper, that's 1 billion tokens. You ingest it in a matter of minutes given enough hardware on this architecture. And with that in hand, you can start

actually making sense of large data sets in the way you would expect of reasoning models. Again, maybe for the developer audience out there, I'm sure you're familiar with use case of AI assisted coding in general, and this is perhaps for currently the frontier use case.

35:31 We are looking at the next generation of use cases like this, but to focus on this use case for a moment, the complexity of having an AI code assistant increases with the amount of preexisting code with a size of the code base. And usually it's much easier to have a model which contributes a piece of new code that just invents things without actually having internalized everything that was created before its action. So it is basically doing a project on the side of its own then to have basically a model which is able to control and contextually operate in an environment which requires understanding of a large code base. And again, code bases are perhaps be the frontier example, but they're still the easiest kind of example that we are looking towards.

Jon Krohn: 36:29 And my final clip from the month of October is from episode number 932 with Larissa Schneider. Larissa recently raised \$50 million through her AI driven company. On frame here I ask how Larissa achieved so much success with the business model that she herself admits, did everything against the book. Tell us a bit more about On Frame because it's a business model that I don't think I've seen before and it seems like it's working really well for you. You recently raised, well, I guess the total of the raises, the venture capital raises you've done so far comes out to \$50 million, including I think a relatively recent announcement. You can correct me on these timings and exact numbers, but this unique model that you have seems to be working out for you. So fill us in on what it's



Larissa Schnei...: 37:15

Yeah, sounds good. Yeah, we started the company in roughly March last year, raised a seed round, then raised another round, so a round in March this year. And I think from day one, we actually did everything against the book. So really not following the typical playbook. If I go back to the very first VC pitches we did, and we came up with this crazy business model and everyone's like, but you need to focus. You can't start with doing something for multiple personas and multiple industries and multiple products from day one. And we said challenges we can, because with ai it's everything has been reset. We were rethinking everything that we've done the same way forever, and we're doing it again and we're doing it better and more efficient and we are really pushing the boundaries in that regard. So when we came out with our out of stealth announcement at the beginning of April this year, we actually came out with, we call it a managed AI delivery platform.

38:14

And in very simple terms, we often actually refer to this metaphor of Lego bricks. So we build an AI platform that is made up of hundreds of different building blocks. So we looked at all of the most complex, the most challenging, the most time consuming problems that enterprise leaders face when building and deploying AI solutions. We packaged it and we use it hundreds of times over for all kinds of enterprise use cases. So someone gets something that is super tailored to their specific environment without having to prepay or no commitment, no cost involved until they actually feel business value. And that's what we came out with from day one. And yeah, it's been working well,

Jon Krohn: 38:59

No commitment, no cost involved until they feel business value.

Larissa Schnei...: 39:04

Yeah, absolutely. That's how confident we feel about it. And it's funny, sometimes people are you sure A POC is

no cost? We're like, yes, it really is not because that's how we build the business and that's how efficient we made the platform. And it's really in Tech On Frame seems to be the only one doing it like that. And the comparable that Chime, my co-founder always imagines, it's like imagine you are getting a new home and you want a sofa, right? Like your custom sofa that fits your specific space and your style and your angles and whatnot, your measurements. Well try to find a sofa build that says, sure, I'll build it for you, totally custom to your measurements and then you can try it and if you like it, you'll pay me. Otherwise, no problem, I'll take it back for free. You won't find that, but at timeframe you can.

Jon Krohn: 39:50

That is wild. And so then how do you know that they're not getting business value and not telling you?

Larissa Schnei...: 39:57

Well, yeah, I mean that's always a challenging area I would say, because what we've seen a lot in AI specifically now it's there's been so much board level pressure, so much executive visibility on the topic of AI that a lot of people are like, let's just execute on it. What can we do? What can we build? Let's just do something. And what we are really pushing for is for them to start with the ROI and the KPIs in mind. So what are you actually trying to achieve? Not just which tech do you have at your fingertips that you could use? And so we really work, we call it a business impact analysis that we do with the customers upfront and say, we want to build one or two or three different POCs with you, but let's try to find the one that actually moves the needle and moving the needle for you means X. And if we hit that, then let's move to licensing.

Jon Krohn: 40:48

I see. I see. So you're kind of with them from the beginning on some metric that they're looking to hit with this particular feature or aspect of their product, their platform. And so it sounds like, correct me if I'm getting

this completely wrong, but it sounds like On Frame is kind of mixing both services and SaaS together. It sounds like you're able to have lots of different ai, AI platform options for your clients that are kind of ready to go, but then you customize them. So there's some services, some adaptation to make the couch say fit perfectly into their space, be exactly the color and the fabric that they want. Okay, so it is a blend.

Larissa Schnei...: 41:38

It is a blend because we think that is very important right now because we've moved so far beyond this moment of generic software. It's like one size fits non, and so we really want to make sure that we offer that, but we don't charge for it. So all of our services and our AI product leaders that work on the specific tailoring of the solution, everything is included in our subscription. So you don't have any hidden costs, no additional charges that just pop up that you never planned on having. And

Jon Krohn: 42:08

Now the subscription that's got to be also bespoke. Presumably some of your clients are using lots of functionality they might add over time. A big client of yours might have lots of different pieces of functionality within their enterprise that depend on you. And so presumably there's different tiers of subscription.

Larissa Schnei...: 42:28

Yeah, we do, yes, but we try to make it as simple as possible as well. It's really fast. It's all about simplicity. We do t-shirt size pricing, so depending on the complexity of your use case, small, medium, large, extra large. But yeah, we do it per solution per year. And some of our customers, as you say, they started maybe with one or two use cases, but now they realized how important on Frame is for their strategy and now we've moved to like 5, 6, 15 different type of solutions that they're running on frame at this stage, but they know how much they'll be paying.



Jon Krohn: 43:03 All right. That's it for today's, in case you missed 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.