

SM287 | 6.13.2026

How to Raise Your Agent | Episode 5

Daniel Solove, Bernard Professor of Intellectual Property & Technology Law at George Washington University Law School

On this week's installment of How to Raise Your Agent, we welcome Daniel Solove into the SmarterMarkets™ studio. Daniel is the Bernard Professor of Intellectual Property & Technology Law at George Washington University Law School. David Greely sits down with Daniel to discuss how AI isn't so much opening up new problems in privacy, but exposing and amplifying the old ones – and how we need to change our legal approach to privacy to solve these problems and take us out of the digital fishbowl we find ourselves in.

Dan Solove (00s):

Go slower, really vet things and test things and do things more rigorously. Don't use it as a shortcut. Take the time to make sure it's working right. Tempting because people want to use the AI for convenience and as a shortcut. The problem is that you don't want to take that detour and then wind up driving off a cliff.

Announcer (25s):

Welcome to SmarterMarkets, a weekly podcast featuring the icons and entrepreneurs of technology, commodities, and finance ranting on the inadequacies of our systems and riffing on ideas for how to solve them. Together we examine the questions: are we facing a crisis of information or a crisis of trust, and will building Smarter Markets be the antidote?

This episode is brought to you in part by Abaxx Technologies. AI agents are fundamentally changing how people and organizations interact with computers and the internet. A race is underway to establish the path that will govern the nature of these interactions. Abaxx Technologies introduces Abaxx Labs as its center for engaging with the developer community through the release of open source software to promote a path that empowers users to maintain control of their identity, data, and agents in digital spaces. Abaxx Labs' first release is the Agents++ library, a subset of Abaxx's ID++ software development kit that has been tuned for use with AI agents. Download Agents++ on GitHub today at github.com/abaxxlabs. That's github.com/abaxxlabs.

David Greely (01m 46s):

Welcome back to How to Raise Your Agent on Smarter Markets. I am Dave Greely, Chief Economist at Abaxx Technologies. Our guest today is Daniel Solove, Bernard Professor of Intellectual Property and Technology Law at the George Washington University Law School. We will be discussing how AI isn't so much opening up new problems in privacy, but exposing and amplifying the old ones and how we need to change our legal approach to privacy to solve these problems and take us out of the digital fishbowl we find ourselves in. Hello, Dan. Welcome to SmarterMarkets.

Dan Solove (02m 21s):

Thanks so much for having me.

David Greely (02m 24s):

Well, really glad to have you here. You know, I am so glad Michelle was able to connect us. I was thinking about where to start today. And one of the things that I keep hearing from our guests on this series is that AI and AI agents are really exacerbating existing problems within organizations and shining a light on many of those. And when I was talking with Michelle about it, she was telling me, you have spent 25 years studying privacy and technology and she has heard you say that AI doesn't create entirely new problems so much as it remixes existing ones in complex ways. I wanted to start off by asking you, what sort of remix of existing problems are you seeing?

Dan Solove (03m 14s):

Yeah, so I think that when I came on the scene, there was everyone pronouncing like it's incredibly new and everything is different and we knew new laws and it's a whole new world. And actually, it's an old world. AI has been around for a very long time. The technologies behind AI are rebranded older technologies that are working a lot better because there is a lot more data and there is a lot more compute. And they are working in ways that are essentially not opening up new problems as much as they are transforming old problems, making them worse by adding speed, adding scale, adding size, and various other things to these problems and mixing

different problems together in different ways. So there is a remix, there is amplification. There are a few new problems that are emerging, but a lot of them I trace back going back a while. And I think it was important to point this out in the article I wrote because legislatures and policymakers were rushing to create new AI laws. Which I think I am not necessarily opposed to, but I think that what was getting lost in the mix was that we need to think about the existing laws and how to retrofit those laws to address AI. In certain ways, the problems with privacy law that I had been complaining about for more than 20 years, actually more than 25 years. These problems, AI really demonstrate emphatically why the laws don't work and why they didn't work. So in essence, I think instead of turning to a new law, why not look back at the existing laws which are trying to address a lot of these problems and failing. And maybe we retrofit and fix those laws to better address the problems. And so that was the general gist of my article, Artificial Intelligence and Privacy. That we shouldn't neglect the opportunity that AI is presenting us to rethink existing law and see that it just doesn't get the job done and AI is really making it crystal clear.

David Greely (05m 45s):

And what are some of those top problems that you are seeing in privacy law?

Dan Solove (05m 50s):

Well, a lot of the law is structured on a flawed model of privacy self-management. That the idea is to empower the consumer, and there is nothing wrong with it, it seems like a noble goal, to make choices and risk decisions about when to share their data and what uses of that data they should say yes to or no to. The problem with this is that this is just too much for a consumer to do. It's too complicated, there is too many companies gathering too much data, and it's almost impossible to make the risk calculation and the reason why AI really puts this problem on steroids is that AI is able to do things that are quite unexpected and it's almost impossible for people to do any meaningful risk calculation when they share data. And so, especially if you consider the way AI can look for patterns in data and can take data and make inferences about other data. So the whole model of the way the privacy law works is that I get to decide what information I am going to share with a company, and then the company has a policy about how they are going to use that. The problem is, if I share even innocuous information, let's say I just share what I bought at a store and I think, okay, this is not a big deal, I am just sharing that I bought hand soap. Well, AI is actually able to make inferences about a whole host of other kinds of information, from political views, to philosophical beliefs, to religion, to health from this data. And so, there is no way that any consumer could possibly know what data they are actually disclosing. When whatever they disclose could be used to make inferences about all sorts of other things they don't want to disclose.

So the real famous case was with Target and this occurred a long time ago, when these algorithms were far less sophisticated than they are now. So Target wanted to figure out what customers were pregnant and this is before they were going to buy baby products. They wanted to figure out this early, so they could start marketing to them and they created an algorithm to look for patterns and they found out that, with pretty good accuracy, who was pregnant and who was not and some of the things that the algorithm spotted were that people bought unscented hand soap, they bought cotton balls. And so if you think about what a consumer is thinking about at the time, they are asked, like, hey, do you want to share your purchases with Target? Sure, no problem. I am not sharing anything sensitive. I am just sharing cotton balls and hand soap, so big deal, right? I don't see a risk there, but now if they knew wow, you are going to share that you're pregnant, that changes the whole equation and this we are now seeing happen in a much greater and more widespread fashion.

So how is it possible for anyone to possibly decide or assess what the risk is when they share data unless they have access to the algorithms and even having access to the algorithms isn't enough. They need the training data, which is impossible for them to provide and they have to essentially run their own tests of the training data into the algorithm to see what the algorithm is going to do. So I think given all this, I just don't think people can possibly make these determinations. If you are asking them, okay, should I share my information? They are making a risk calculation. They are looking at, okay, what's the benefit of using this technology or sharing the data? Am I going to get a discount? It's going to help me use this app and then they assess the risk and they make a decision. Are the benefits better than the risks? That's the way the law is trying to set that up to empower the consumer. But if that becomes meaningless, then the laws are meaningless. Then what are they doing? What's the purpose? If I can't make that risk calculation, then it's just a game of gotcha and all sorts of unexpected things happening and the law is completely ineffective and it's creating the illusion that people are consenting to all this and agreeing to all this. So when the gotcha comes around and something's used in a way that I don't like or that's harmful to me, it's like, well, you agreed. You gave the data. You said, yes, you clicked this box or you didn't click this box or you visited the site and you didn't opt out and that's all fiction. We are not really agreeing to this. We don't know. People don't know what's going on. Even as a privacy expert, I don't really know what's going on.

So I think all of this really shows that the law's current approach doesn't work. It never really worked. It's not effective. We need something else, something that really will work more effectively. And I think AI makes it so clear because it is really complicated and it's doing things with data that humans can't do. Humans can't process billions of data points to make these kinds of inferences and see the kinds of patterns that AI is able to see. That's the kind of special magic that AI can do, is it can see patterns that humans can't see. And so to say that humans are somehow supposed to see the pattern is really contradictory and we need to really rethink what we're doing with privacy law so that we can protect people and ultimately, I don't think, given how complicated things are, that the average person is going to be able to make these risk calculations in any meaningful way. So the law needs to protect them. Just like when I buy a car or I buy a product, the law has my back and I know that, hey, the car's not going to blow up on me. And if it does, I can sue. There is going to be hell to pay, but someone has kind of made sure it's safe. When I go on a plane, I don't have to be an expert on flying a plane. I don't have to fly it myself. I don't have to know airplane mechanics. I go and buy milk at the store. I don't have to become an expert on pasteurization. I don't have to research every farm to see if the milk's going to be safe. I know it's going to be okay. So many things, I go through the day, and we don't have to become experts on it. We just know that, you know, someone has our back and that these things are not going to harm us and we don't have to worry about buyer beware. The problem with privacy and AI these days is that we do all the onus is put onto the consumer to somehow figure it all out.

David Greely (12m 39s):

It makes a lot of sense of needing to move from this existing mental model of a super rational consumer who can take in all this information and make the risk assessment and provide consent to one, which is, you said, more like purchasing a car or when we take prescription medication, we have some knowledge that it's safe and effective and it's been tested and when you think of that in terms of privacy, where would you think the regulatory bounds would be set in terms of the use of data being safe and effective, or at least safe? How do you think about how much privacy is the right amount?

Dan Solove (13m 19s):

Well, I think we need both an ex ante and ex post type of component to the law. So I am not a fan of strong government paternalism and the government just telling me, okay, this is what I can do and this is what I can't do. I don't want that. On the other hand, I do think we need adequate safety testing of various things for privacy and all the harms that it could create. So AI models and so on. We need to make sure that when they are put out on the market, given to people and uses of it are pushed upon them and I think the companies are trying to push every which use possible, like throwing everything up against the wall and seeing what sticks in this massive attempt to see what might be profitable for AI. But when I buy a car, there are certainly a lot of choices about safety that I can make. I can buy a safer car and a less safer car and there are all sorts of things I can read that will explain to me what is safer and less safe in a car. But I know like I am not going to be able to buy a car without seatbelts or airbag. I know that certainly the car is not going to blow up on me like a four pinto. So I know that someone's kind of looked at it in advance doesn't mean that I don't have choices, that I can't buy a safer car or a less safe one. I still have choices and the companies can still make a profit and can still make cars that are more or less safe but there is at least a minimum and there is also, I know that if something goes really wrong, that there are ex-post remedies. So if something goes wrong and I am harmed, a car malfunctions, I can sue and the companies know that and as a result, I think we have cars that are much safer than they were before.

When the industry was furious, when advocates were pushing for more car safety and saying cars were not safe, they were outraged. They said, you know, people don't want safer cars. How dare the government possibly regulate this? It will kill the car industry. We won't be able to have cars. We won't drive. Everyone's going to just walk. Again, they will ride bikes and walk. But of course, that's not what happened. In fact, we have safer cars and it's interesting that one thing I think that companies are actually really good at doing is they respond really well to incentives. You give them the right incentive, they respond and so if we say we want you to make safer AI and we create the right incentives to do it, they are going to do it. We say, hey, we want cars to be more fuel efficient. We have more fuel efficient cars. We want airplanes to be safer and they are safer and that's the thing. Generally, when we have demanded things like, hey, we want products that aren't flammable, we get products that aren't flammable. We want food that's safer. We get safer food. And I don't understand why these lessons can't be learned in technology, that we demand it's going to be safer and if there is a mishap, yeah, there is going to be hell to pay. You are going to get sued. There are going to be consequences. They will put more money and effort into making it safer and less harmful and unfortunately with tech, policymakers are under the spell that anytime you say the word innovation, that they will go, oh my gosh, let's just drop everything and do that. As if you have said like some secret password that unlocks, they will just drop it all and fortunately, that's the lessons that we saw that worked for food, that worked for drugs, that worked for cars, that worked for planes, that worked for almost every single other thing. I don't know why we can't take those lessons and put them into technology.

David Greely (17m 13s):

And I wanted to ask you about how we should think about the harm done by not respecting privacy. I think so many of us have become accustomed to clicking through and agreeing to things because we want the convenience. We are certainly an attention-seeking society right now. Do you have some practical examples or thoughts on the harm being caused by not enough privacy being respected?

Dan Solove (17m 45s):

Yeah, so I think we have way too much data being collected and used in ways that people were not anticipating and in ways that could harm them. So all sorts of data that now could be obtained by the government with the most minimal procedure at all. So it can be dogged and grabbed up, used for enforcement of pretty much anything, whether it be by ICE or by a state that wants to criminalize abortion or by a state that wants to criminalize anything, or go after people for what they are saying or doing or whatever. It could be a lot of different things where the government could easily get this data from a company that has gathered it and I think a lot of people don't realize that so much of their data now is in the hands of companies and there's hardly any protection from the government obtaining it and even when there are strong protections, like the requirement of a warrant, which is often not required, but when it is, a warrant's easy to get. It's basically a rubber stamp and so there is one case written about in a great book that my colleague Andrew Guthrie Ferguson wrote this great book called *Your Data Can Be Used Against You*. He is a criminal procedure scholar and he writes about this really fascinating case about this guy's pacemaker and the pacemaker collected data about his heart rate and that data was used against him for an insurance fraud case because it contradicted his story about what happened. So your body can be used against you and I don't think that people, when they get these devices installed and they are increasingly putting these devices into our bodies and wearing them and carrying them around with us all the time, that they could, in fact, be used in ways that could put us in jail and do a whole host of other things, be used to deny us insurance, for example. Hell yeah, we just want to track your driving or we would want to track all your supermarket purchases, but what if an insurance company decides, hey, we think you are buying too many French fries and too many fatty foods. So we don't want to insure you because you have a history of heart disease and we don't like how your heart is beating or your data from your watch, which shows a high heart rate and so on and so forth. I mean, there's a lot of existing cases that are happening and possible cases that could happen and I think that the law just is not addressing any of it and essentially, the law is you are on your own. If you don't like it, then you shouldn't have used the product or service or the AI Chatbot, you shouldn't have done it. You agreed by using it and you are tough luck. I just think that approach is too much of a gotcha approach where somewhere down the road, something is used in a way that you don't want and expect, and then you are screwed.

David Greely (20m 57s):

And so many of these things, the idea of consent when you really can't be part of modern society without using a lot of these tools. So the idea that you can opt out is a fiction as well, isn't it?

Dan Solove (21m 12s):

Yeah, I mean, you can't live really without a phone anymore and so now we have this device that's tracking where we are going all the time and collecting massive amounts of data. A lot of times you can't do much without doing stuff online and transactions that you can't do without being tracked. I mean, it used to be that you could do a lot of things in relative anonymity. You go to the store, you pay cash for something, your data isn't tracked but now it's possible to have this enormous stream of data about you with every single thing you do, everywhere you walk, it's hard to escape from being on a camera and then facial recognition on top of that and so the idea that you can kind of go through life and most of what you do is not recorded and not tracked in some way and there isn't this massive repository of data about you, that is really a thing of the past and it's almost impossible to live in today's age without that and no matter what you do, it doesn't matter in a way. People will often ask me, like, what should I do to protect against fraud, identity theft? How should I guard my social security number and I can tell them, you know what, you can take your social security number card and you can lock it in a safe and then put that safe within another safe and you cannot give it out at all and you can keep it so super safe and you know what, guess what? Any company can sell your social security number and they are already selling it online. So anyone can just buy it. Anyone can find it. It's readily available in public record. So you can do everything right and it will not help you at all and that's typically the case with a lot of these things, that people can take all these steps and it creates the illusion that they have some kind of control. But they really don't at the end of the day and most of the things that people are being asked to do are kind of pointless things, like rearranging deck chairs on the Titanic that don't really change the equation and don't really, yeah, and marginally make them safer, but not that much.

David Greely (23m 20s):

And one of the things that it seems that the use of AI agents, one of the problems it's accelerating, is now to go back to your first example with Target, and we have companies that are collecting data and the AI, the agents, lets them process it much more effectively

and understand more about us from it. And then you have got these companies potentially acting as agents for the government. But now we have our own agents who can, in essence, be like the mole in the house, revealing information about us, leaking data out, or even potentially consenting to data releases on our behalf. How big of an issue do you see that becoming?

Dan Solove (24m 05s):

I think it would be a very big issue, and it kind of depends on what the AI is being used for and what I see is it's being used for almost everything, from like a personal butler, to a companion, to a psychotherapist, to a lawyer. So people are, oh, I am going to use my Chatbot for legal advice, and start asking and talking to it like a lawyer and then they are shocked to find out there is no attorney-client privilege, and that everything they have done in the Chatbot is now being used against them by the government, whereas if they went to a lawyer, it's different. Or people using Chatbots to engage in psychotherapy, but these bots are not trained like a psychotherapist. They don't have a license like a therapist. They haven't got nothing to lose and then the Chatbots are being overly affirming to delusional thoughts, are being overly sycophantic, and not pushing back enough on somebody, or not appropriately recognizing when someone is talking about self-harm or harm of others. They are easily circumvented when people say, oh, I am just asking for a friend, or I am just writing a book, and I need it for my fictional story and they can easily evade it. All this interaction with the bot, I mean, people think that they are talking often to a person and even when they know that it's AI, the simulation of a human on the other side, the bot talking as if it has a body, as if it's a person, that, even when people know it, they still treat it and are lured by the human simulation to a point that the knowledge isn't going to cure the problem. They are more forthcoming. The movie *Her* really captures this well. It's a movie made about 10 years ago. It's a wonderful movie about this guy who falls in love with this AI on his phone and he knows it's AI, but he doesn't really clock it. It doesn't really sink in. He doesn't fully appreciate the implications of it until later on in the movie. I don't want to spoil the movie, but there is a scene where he really realizes, like, oh my gosh, this is not human or even human-like, but the simulation there gets people to say things and talk about things that are deeply personal and all this information, it's not confidential, like it would be in a relationship with a therapist. This is information that is gathered by a company and it's data that a company gathers that is similar to any other data that it's gathering about you and can be accessed by the government or used in every which way. That's what's going on but I think people are lulled into this false sense of intimacy with AI that gets them to reveal a lot more information than they would possibly reveal if they were given a form on a website or a company's representative were to call them up and say, hey, hi, I am from Facebook. Will you tell me your deepest secrets?

David Greely (27m 23s):

And I think you have written on this recently, but the idea of humans hiding behind agents as being the one that should be accountable for doing highly unethical, if not illegal, things. Kind of the agent did it, the model did it, the algorithm did it. When you hired somebody to use your last example to pose as a psychotherapist and pretend and then go and extract information, I imagine there would be some liability in doing that, but creating an agent that does that, at least so far, seems fine legally. How do you think about how we hold people, like, where does the legal liability fall in a world where we are increasingly using these AI agents?

Dan Solove (28m 13s):

It's an interesting, great question, and really interesting. I think we have two layers here, two possibilities. One is the companies that make them and how they market them and then there are the individual users and I think we see, in both cases, the attempt to kind of shirk responsibility or avoid responsibility. So the users of it, like, I use a Chatbot and I have a Chatbot do something for me. I think I should be held responsible for what the bot does. So if I say, hey, I am going to have the bot write something for me, I am responsible for what that bot says and I think a lot of lawyers are using it to write briefs and then the AI is making up case sites, and now they are getting in trouble and they go to the court and they say, oh, I didn't write it, it's all the bot's fault. Well, you shouldn't be using the bot for that. The tool is built to make things up. That's what it is built for and so it's going to make up cases. That's what it does and there is attempts to try to make it, try not to do it, but essentially it's essence, I think what it's generally designed to do, what the technology generally does, is make stuff up. That's what it's up to. You can try to fix that and correct for that, but ultimately you are using a technology that's being retrofitted for something that its original structure is different. So using it for that purpose, someone should be responsible for it. You use AI. If I hired a person to do something, I am responsible for that person. A company can't hire a person and then say, oh yeah, they are my employee, but they did something wrong, like we are not responsible. You are responsible for what they do. Whether that person's a human, or whether it's a bot and I think there is an attempt to say, well, it's AI, it did stuff I didn't expect and they want to use the AI as a shortcut. The whole point is that you want convenience, you want the AI to do things that don't want a human to take the time to do, but you still have to review the output to make sure it's accurate and that review is tricky and it's tricky because the AI is great at simulation. It's great at writing things that sound good on a quick read. It sounds pretty plausible but then when you really read it closely, you see it's kind of hollow, it doesn't always add up, but it's great on the quick read because that's what the technology's good at doing. I blew booking sites, which is getting a site into the proper format. It's annoying, it's tedious. So I

thought, I want to use AI for that. Wouldn't that be great? So I said, okay, here is a newspaper article. I gave it the URL of a newspaper article and said, okay, can you put this into a citation? Pull all the information from the article and give me a citation to the article. I am like, look, I am not even trying to have it interpret the article or write anything about the article, just do a site, right? Simple as that and this is like a New York Times article and it gets the authors wrong. The problem is it's not even doing the things that are easy, particularly well, in certain cases. It can do other things but when you looked at it, it looked pretty good. Like in one case, it was like there was a New York Times author. It was a writer for the New York Times, it just wasn't the writer of that article. So it kind of looked plausible. Everything looked generally fine but then when you really took a close look, you saw, oh, wait a second, they weren't there. It was a different article, but why would they get it wrong?

David Greely (31m 49s):

Yeah, sometimes I wonder how much it plays into our confirmation bias, that it gives us what we expected to see and so on that first pass, you are like, oh, yeah, that looks right, until you dig into it a little deeper.

Dan Solove (32m 01s):

And it looks good. Like it generally, I mean, except for the instances where you get a picture of six fingers or something, a lot of times what it generates when it writes, it just generally sounds good but then if you really pick it apart, if you really look closely, that's where you see it messes up, or it's mostly good, but then there is a hallucination in the middle and that's what makes it so tricky to find. If it were, everything were wrong, and it wrote terribly, it would be less alluring but then there's also the problem of the, that's the user of it and their responsibility, but then there is the responsibility of the companies because the companies are putting out this tool and they are desperate to make AI profitable and to find uses for this thing. So they built this, kind of this everything tool and then they basically are throwing it out to the public and at every turn, everywhere I go, there is someone in my face saying, like, user AI, user AI.

It actually reminds me of the book, Green Eggs and Ham, where Dr. Seuss, where there has that annoying guy, he is like, do you want green eggs and ham? Green eggs and Ham do you want it with a cat? Do you want it with a bat? Do you want it with a car? Do you want it with whatever? And I feel like AI is the same thing. I can't go for more than one second with, like, do you want AI for this? Do you want AI to write your document for you? Do you want it to, do you want to use it to blow your nose on it or whatever. Like, whatever it is, the AI is constantly being shoved into my face for every single use and they are pushing it. Like, hey, use it for medical stuff. Use it for this. Use it as your companion and then the thing is people use it and they often misuse it. They don't use it well. It doesn't turn out so well and then the companies say, hey, not us. We didn't do it. We just made the tool and the person didn't use it well. It's kind of like giving someone something that is like a firecracker and say, use it, whatever. Put it in the oven. See what happens. You put it in the oven, it blows up. And then they say, oh, it's not our fault you know, why would you put it in the oven? Well, because you have been touting it as the next exciting thing and you should be using it for everything. I think that's part of the problem.

If the companies are actively pushing this thing for every conceivable use, without testing for whether it's appropriate for that use or optimized for that use, and then trying to disavow responsibility when people start using it for irresponsibly or for bad things, that's a problem and I feel like it kind of reminds me of a little bit of the Q-tip. You shouldn't stick them in your ear but that's what people buy them for and so with a lot of things, I think the law has evolved to a point that if you put out a product and it's a kind of known thing that people might use in certain ways, you need to effectively provide protections against that so that it doesn't harm people. Drug manufacturers know that you might tell people take only one of these pills a day and they know that people will probably maybe take two. So it shouldn't be that if you take two, it kills you. You should have some degree of, now obviously you can't do that for everything. People shouldn't drink Drano. But there is a lot of things where we have forgiveness in the margins. We have protections. We know that things are going to be used in bad ways. We wouldn't say, hey, you know what, you are only supposed to drive the car the speed limit. So if it's above 65 miles an hour, the car's going to blow up if you have a crash. We say, no, we know people are going to speed, so we have to make sure the car is safe even when they speed. That hasn't happened with a lot of the digital technologies and AI, and I think it needs to happen but there is such a rush to kind of get this out and have people try it for every single conceivable thing, some of which it's good for, some of which it's not good for, some of which it could be good for if there were guardrails and protections built in and the tool where time was spent to optimize the tool for a particular use but there is just this rush to like, let's just throw everything against the wall and we are going to have to find a use that's going to make us a profit. So let's just throw it all out there and see what happens. That, to me, is not the wisest approach.

David Greely (36m 26s):

You are teaching the next generation of lawyers at the George Washington University Law School. Through Teach Privacy, you are working with many Fortune 500 companies and other large organizations and I was wondering, what conversations are you having

there? What are you trying to teach or what are you seeing organizations and the legal profession doing that you would say, okay, I think this is the right way to approach AI in this moment?

Dan Solove (37m 01s):

I think it's hard because there are so many different ways that are occurring. So you teach a student, like, use AI this way, but then they go to a firm or play a business, and they use it in a different way. So there hasn't really been an appropriate kind of set of guidelines and approaches for the use of AI in law, for example, or the use of AI in a lot of things. It just hasn't been developed yet, and there is a lot of dispute over how much it gets used and what the appropriate ways of using it are. So I think we are really in a place of the Wild West at this point where it's hard to know what to say. I generally give the advice of be very cautious, that we are dealing with something, an exciting tool that I think has tremendous potential. It can really do some amazingly great things but I think that a lot of the things it can do need human hand-holding and human involvement, and that isn't quite there yet.

The protocols and the ways that we can ensure that they are working properly are not built yet, not developed yet. So you are dealing with something without guardrails, and they are not there, both built into the product or imposed outside of the product by law or even appropriate guidance and so it's a very, and we don't fully know all the harms yet, and it's very much a dangerous tool to use because it can have a lot of negative impacts. It can have a negative impact on a student who uses it as a shortcut to do writing or thinking even. I would say that certainly you can't use it to do your writing for you, but also even if you use it as an aid to help you study, that could be a problem. One, because it could get things wrong, two, because it could also take a shortcut around thinking time and that's one thing that's tricky, is it takes a long time in society for us to figure out what types of things we can jettison and what the costs are. So an example in law is Shepherds. They used to be you wanted to find out if a case was overturned or what happened to a case later on, how many times was it cited and you had to go to a book called Shepherds, and it was a ridiculously onerous process of looking through the book to see fine citations and where it was cited and everything else and technology took that away.

Technology now, you can track how a case is cited very, very easily without having to go through those books and no one uses Shepherds anymore and you know what? I used Shepherds back in the day, early on in my career, and I don't miss it. I really don't think there is a big loss in making this shift with technology. I have lost nothing. The time I spent shepherding did not give me greater wisdom or help me in any way other than just take a lot of time. So great, that's a win, easy win. The problem, though, when AI is doing analysis and research and writing, it is having an effect. It's shortchanging thinking and what we don't fully know yet is how much value that thinking has. Sometimes it could just be tedium, like if it's just correcting citations, I think that's tedium. Who cares if someone knows a citation format? I don't see a benefit. But thinking through something, researching something, has a lot of value. I remember, like, you go to the library or go to a bookstore, and you are walking around, and you just stumble upon a cool book that you didn't know. How great is that experience? How many times as you are reading something or researching something, you find something interesting, or you have an idea because you have been steeped in something, thinking about it for a long time. I think those are things that we lose when we delegate everything to AI and I don't think we should lose them. They are not necessarily so efficient, but there are a lot of wonderful gains and wonderful things that happen. And if we don't let people have their brains think and process, their studies are starting to show that people are less smart, and they become less good at doing those kind of analytical tasks and so we don't want AI to systematically start weakening people's writing ability and their thinking ability.

We want to find ways that AI can take over tedious tasks that ultimately are not a productive use of time, like citations. Sadly, when I tried to use AI to correct citations or make citations that were of actual sources, it didn't quite get the job done, which was sad for me, because that's where I really wanted to use AI, is in these kind of more thoughtless tasks, where I don't want it to think for me, I want it to do tedious work for me that the non-thinking work for me is where I really want its help, and it wasn't helpful.

David Greely (42m 27s):

And you have mentioned that there is some fundamental issues, some fundamental problems in privacy law are being exacerbated by AI, that the legal frameworks for the use of artificial intelligence don't really exist yet, and there's a lot of ambiguity around this. And I was wondering, if you were advising a company that's about to deploy its first generation of AI agents into some consequential part of its business, what would you want to feel confident that they have thought through?

Dan Solove (43m 05s):

First of all, I would tell them to be careful about the hype. The AI companies have a very strong incentive to overhype it. They are desperate to find especially enterprise users of their AI, they are desperate to build business and basically keep the lights on, which are kept on. AI is very expensive, and so it takes a lot of money and resources to do it, and a lot of investment money that eventually will

dry up if it doesn't ultimately turn a profit. But it's not turning a profit now. So I think that don't buy the hype. So be very skeptical about all the claims, because it's probably over claimed. That doesn't mean it's not valuable and it can't do good things, but the problem is that AI companies are kind of selling it like a used car salesman and then I think really make sure that it should be used in ways, I think, to find ways to augment human productivity rather than replace it and if the idea is like we are going to do a shortcut, how can we replace humans and put this AI in there? I think that's typically where things go awry because the AI makes a mistake, and it doesn't work as planned, and now you are kind of screwed. Because you put a lot of weight and trust into the tool that is not going to perform well and now you have lost your human teams that can do that. If you can find ways the AI can work with people and make sure it's tested out and works well and is doing what it does and performing well, then great. But make sure it's tested and it takes time, and it takes a close look.

Because one thing I have discovered is that AI can do things and make them generally look good enough on a quick glance, but then the closer look reveals its shortcoming and its problem and we have seen this time and again where people think it looks good enough, and they go out with it, and then it winds up coming back to bite them and that's because they got lulled into thinking it seemed good. And so we have to make sure it really is good. And that means more closely monitoring what it's doing. So go slower. Really vet things and test things and do things more rigorously. Don't use it as a shortcut. Take the time to make sure it's working right. It's tempting because people want to use the AI for convenience and as a shortcut. The problem is that you don't want to take that detour and then wind up driving up a cliff, which is often the case because people who want a shortcut will take a shortcut to get to the shortcut and I say don't take the shortcut to use the shortcut. Really make sure you are getting a quality output from what you are doing with the AI that's working right. Because if I had a nickel for every case I have heard about or instance I have heard about where AI seemingly does things okay and then I found out it really screwed things up, it can be bad. Just an example, I was talking to someone at a conference the other day and they told me about how they had the AI, wanted the AI to take some documents in Word and turn them into PDFs. But it actually took them and replaced existing PDFs that this person had where they had some annotations in their PDF and it overwrote those annotations and then said, well, can you recover it? Sorry, I can't help you there. So just one little example. But I think we see many examples of that happen where people want a shortcut, seemingly seems fine, and then you try it and then it winds up being, creating a mess. So don't shortcut into your shortcuts.

David Greely (47m 04s):

Well, thank you so much, Dan. Really appreciate you taking the time to have this conversation with us and I wanted to ask you one more question before you go. And that is, you have written a children's book about privacy. What led you to want to write that book and what would you like the younger generation, how would you like their attitude towards privacy to potentially be different from the current generations?

Dan Solove (47m 32s):

So I wrote my book, it's called *The Eye Monger*. I wrote it because I would always have story time with my son and we would read books, but also I would tell stories and put him as a character. He would often have influence in the stories, you know, telling me to steer it one way or the other and I realized that as I read all these children's books, that none of them talked about privacy and I thought there really should be more talking about privacy and technology and the role it plays in our lives. This is more than just privacy, it's AI, technology, everything. There is hardly any books, especially fiction books, that involve the use of technology in our lives. And so I thought this would be a good book to create and I have my son as a character in the book, just like our story time. I did it with an illustrator that I use for my privacy training company. He does wonderful illustrations for me for my courses, and I make these whiteboards, which are these one-page visual summaries of the law, and so he illustrates those. He also illustrates my cartoons, and so I thought it would be perfect to illustrate this book and bring it to life. So the book is the story about this guy, this creature who has all these eyes all over him, and he goes to this new town, this town, and the people he says, I am going to watch you and keep you safe and they elect him to do this, and things go awry. I won't spoil the book, but that's sort of how it resolves, but that's the story and the idea is to get young people to start thinking about privacy and why it's valuable, and what are its virtues and even its vices, like why some might want to sacrifice privacy and what the benefits and costs are. So that's the book and why I created it, and I hope that there will be more books on these topics, because I think that young people really do need to be educated about the world they are living in, and this is a world of technology. It's a world where everywhere we go, there is a surveillance camera and facial recognition and devices and data being gathered. It's a very different world than the world I grew up in. I think a lot of people, oh, young people, they don't care about privacy anymore. Look at what they are doing. They are sharing all their data. Well, they don't have a choice. To live in the world today, you have to do a lot online. You know, there is a lot of times where if their school were canceled, there was snow, you stayed home, but now you Zoom. And kids play with other kids online through a video game, but all this stuff involves data being gathered about them and they are living in a world where everyone has a phone, and they can be filmed at any moment. So I don't

think it's fair to say they don't care about privacy, because the world is forcing them to live in an environment where they have much less privacy and I think a lot of them understand that privacy is not because they are choosing. It's because the world has stacked the deck against their having any privacy or not really giving them meaningful choices and so ultimately it's the job of the law to mandate and protect them so that they do have the protection that they need and that their data can't be used in certain ways to harm them, and that they do have these protections because there is a convenient myth that they don't care, and therefore we shouldn't protect their privacy because they've just given it up because they are all texting and they are on social media and they're using the phones and they are posting online and they are doing online video games. So I really do think that we have made the world that they have to inhabit and the world we have made is we have made a fishbowl for them. We don't decry the fish because they are in the fishbowl.

David Greely (51m 39s):

Thanks again to Daniel Solove, Bernard Professor of Intellectual Property and Technology Law at the George Washington University Law School. We hope you enjoyed the episode. We will be back next week with another episode of How to Raise Your Agent. We hope you will join us.

Announcer (51m 55s):

This episode is brought to you in part by Abaxx Technologies. AI agents are fundamentally changing how people and organizations interact with computers and the internet. A race is underway to establish the path that will govern the nature of these interactions. Abaxx Technologies introduces Abaxx Labs as its center for engaging with the developer community through the release of open source software to promote a path that empowers users to maintain control of their identity, data and agents in digital spaces. Abaxx Labs' first release is the Agents++ library, a subset of Abaxx's ID++ software development kit that has been tuned for use with AI agents. Download Agents++ on GitHub today at github.com/abaxx-labs. That's github.com/abaxx-labs.

That concludes this week's episode of SmarterMarkets by Abaxx. For episode transcripts and additional episode information, including research, editorial and video content, please visit smartermarkets.media. Please help more people discover the podcast by leaving a review on Apple Podcast, Spotify, YouTube, or your favorite podcast platform. SmarterMarkets is presented for informational and entertainment purposes only. The information presented on SmarterMarkets should not be construed as investment advice. Always consult a licensed investment professional before making investment decisions. The views and opinions expressed on SmarterMarkets are those of the participants and do not necessarily reflect those of the show's hosts or producer. SmarterMarkets, its hosts, guests, employees, and producer, Abaxx Technologies, shall not be held liable for losses resulting from investment decisions based on informational viewpoints presented on SmarterMarkets. Thank you for listening and please join us again next week.