

Actuaries Institute Podcast

Driven by Data: An expert's take on ethical AI

Date: 15 July 2022

Interviewer: Meg Yang

Guest: Chris Dolman

Duration: 31:09

Meg Yang: Hello, and welcome to the Driven by Data Podcast, an initiative of Actuaries Institute's Data Analytics Practice Committee, also known as DAPC. My name is Meg Yang, and I'm a member of the DAPC's Young Data Analytics Working Group. Joining me on podcast six of Driven by Data today is Chris Dolman. Chris is currently the Director of Data and Algorithmic Ethics at IAG, helping to ensure that modern decision-making algorithms and other advanced users of data are designed and implemented in an ethical, responsible, and thoughtful way. Chris is a Fellow of the Actuaries Institute, a Gradient Institute fellow, and an active member of Actuaries Institute's Data Analytics and General Insurance Practice Committees. Welcome to the podcast, Chris!

Chris Dolman: Thanks, Meg. Hello, everyone. Thanks for having me.

Meg Yang: Let's start with a brief introduction of yourself.

Chris Dolman: Sure. Well, thank you for a nice introduction already, so maybe I'll elaborate a little bit on that. So, as you can probably tell from my accent, I'm not from Australia. I grew up in the UK and worked there after university as an actuarial consultant for a little while, and then took advantage of portability of the actuarial profession, which is nice. Moved over here to Australia and moved out of consultancy and into industry and I've stayed in industry ever since. So, I have done pretty much all of the traditional roles that actuaries can do in GI, so a bit of reserving, a bit in capital, dabbled in portfolio management for a little while, worked in pricing for a little while, did a few other things.

Chris Dolman: And then the last few years I've been heavily involved in the analytics space, particularly outside of the traditional areas of pricing and underwriting. I've gotten quite active in these new areas that we can get involved in. As you mentioned, I'm also a Fellow at Gradient Institute, which is a not-for-profit based in Sydney focusing on responsible AI. So I do a lot of collaboration with them and an awful lot of volunteering for the Actuaries Institute, as you mentioned. So, that keeps life interesting.

Meg Yang: **Yeah. So, you mentioned responsible and ethical AI. How do you define ethical AI or ethics in AI?**

Chris Dolman: Yeah. Well, that's a question that often comes up and it's not one that I particularly like answering in terms of a definition, because although I have ethics in my job title, I find discussion of what it means tends to get in the way of a productive discussion. I mean, I've witnessed heated debates between people arguing about whether we should be worrying about responsible AI or ethical AI or trusted AI or other words that get thrown around. What happens is that people get very passionate about particular words getting used to describe the thing that they want to describe, but quite often they're really talking about the same broad topic.

Chris Dolman: And so, I try and avoid those arguments because I don't find them particularly productive and I just try and focus on just practical action that people are trying to take. So really, in my day job and in my volunteer activities in the profession, we're really just looking at making sure that automated systems or other advanced uses of data, you might call them AI. You might not. Let's not agonise over whether we count it as AI or not because again, that's a bit unproductive.

Meg Yang: **That's a different podcast. That's a whole debate itself.**

Chris Dolman: But again, it's a fairly unproductive debate, right?

Meg Yang: **Indeed.**

Chris Dolman: Does it really matter whether we classify this thing as AI or not? It's this thing, whatever we call it, designed to do something that's appropriately legitimate or justified or socially accepted. And then if that's the case, have we actually built the thing in the way that it was designed, right? Is it actually doing what it was designed to do without horrible side effects? If we can tick both of those boxes, then we can probably be quite happy. We might call that responsible AI. We might call it ethical AI. There's all of these dimensions that you can throw into that.

Chris Dolman: So, just focusing on ethics, there's a bunch of topics that often come up, which you have to think about when you're building these systems. So, there was the decision that you made using your system, just who benefited from it? Who suffered detriment? Is that appropriate? Is that fair? Is that acceptable? Was the decision made for sound reasons? Could you explain those if you had to? Should it be able to be challenged? If so, how does that work? Et cetera, et cetera. So all these sort of topics come up in the debate over ethical AI or responsible AI and many, many more, but I don't agonise too much about the label that we give those topics. I just try and get on and think about them and do them.

Meg Yang: Fair enough. I think it's great that you've defined the broad terms for the ethical AI or ethics in AI for our audience so that they know what we are talking about later on. You had a wealth of experience in consulting industry pricing, underwriting, but what sparked your interest and brought you into the areas of ethical AI and responsible AI initially?

Chris Dolman: Yeah. Well, I used to run the centralised analytics function at IAG, which we had some time back. We had big plans in that area to use more automation and data-driven decisioning. We were really scaling out what we were trying to do in that team. It was pretty clear to me running that team, that there were risks involved. I'd noticed around that time, this was mid, late 2010s, increasing stories of AI gone wrong. This was before Cambridge Analytica, but in that period where AI was getting used more in industry and in government. There was all stories every week about something that had gone wrong, right? So I thought, "Well, there's issues here."

Chris Dolman: I had a look around for written guidance, went and talked to my favorite consultants about the topic to see what they thought and everyone was a little concerned about things that could go wrong, but there was no real substantive guidance on what to do about that problem. There was some guides on model risk management, mostly out of the U.S. because there'd been some regulation there some time before, but it really focused on mechanics of model building mainly, not so much on the ethical questions I mentioned before, sort of intent and legitimacy of what it is you're doing. So, that was a bit of a gap in terms of regulation and guidance that I could see.

Chris Dolman: There were some early AI ethics frameworks that had come out all around the world, but I found those fairly unhelpful. They were pretty vague and pretty high level and super high level principles in most of them, but very little underneath. And so, after doing all that and the news stories kept coming during that time I was exploring the topic, I thought, "Well, there's a need to contribute here." I was also in the mindset of a change of career direction at that point, so I really just decided to try and jump into the space internally in the creation of the role I'm in now and also trying to do more outside of the day job, so with the institute and just more broadly in society, writing papers and presenting them and giving talks and just thinking about the topic.

Chris Dolman: We also set up Gradient Institute around the same time, which I mentioned before. So I do quite a bit with them. Look, it's a huge topic to get into, right? So, as one person, you are never going to completely solve that problem. There's a way too much than one lifetime can give you in terms of time. So there's no chance of getting bored or running out of things to do, which is good. There's a nice blend of technical and nontechnical things to cover, so that keeps someone like me interested, who wants to straddle both of those domains. So yeah, it's been good.

Four years in, there's still plenty to be done, so hence this is one of the jobs I've been in the longest in my whole career actually.

Meg Yang: **Wow. Maybe there's more and more jobs coming up. So, are you-**

Chris Dolman: Well, potentially. Yeah, that's it.

Meg Yang: **Are you able to give some specific examples of the research and publications you've done?**

Chris Dolman: Yeah. Well, there's a lot, so maybe I'll just mention the more recent ones, but if anyone wants to have a look, there's loads on my LinkedIn or we can get in touch and I can send you a few links. So, there's some that are quite theoretical, some that are quite practical. There's a case study or two there as well. So whatever you're looking for, hopefully I've written something in that general vein. So the recent stuff I've done, which got presented at the Actuaries Summit recently, which might be of interest to people, first one's a collaboration with my colleagues at Gradient Institute and this really spells out some risks caused by AI systems. It focuses on governance risks and then identifies just a series of practical steps that you can take to try and reduce those risks.

Chris Dolman: It was something that, again, guidance out there is quite limited, but people need some guidance on what to do today, so we thought we'd write some down. It's been fairly well received, so that's been nice to see. We published that prior to the summit, and then the summit was a good opportunity to promote that to the profession, but it's definitely been one that we've been promoting to other groups as well. It's written to be very accessible for general management, right? So you don't need to be like a data scientist or anything like that to get some value from it. We're really hoping that senior execs can read that and get something from it in terms of how they govern these systems.

Chris Dolman: Second paper I presented at the summit was on Consumer Data Right. So this is a topic I've been getting quite interested in, which is... I mean, when you talk about sort of AI ethics, it's important to focus on the stuff that goes into AI models, which is data, as well as just the models themselves. And so, I've been getting quite focused on data, as well as models in the last little while. Consumer Data Right is a huge reform happening in Australia, where data sets are being opened up that were previously not opened up, which I support as a general concept, but there are risks involved, and so we need to think through quite carefully how to manage those risks.

Chris Dolman: So this paper identifies some particular areas of unfairness that we can see that might emerge in the regime that's currently designed. And so, we try and spell those out so that people can understand them better and we can have a discussion about how we might resolve them. It's something

I've seen as a bit of an issue for a while. So I thought as CDR expands and goes into more sectors and areas of the economy, it's something that we need to, I think, have a discussion on. So I wrote this paper with my colleagues at Sydney Uni, couple of legal academics I know, so we could make these issues clearer to others. So hopefully that, again, has some impact. It's an example of what you can do, I guess, more narrowly than the first paper, which was quite broad. So this area gives you opportunities to contribute both broadly and holistically, as well as in quite a narrow targeted way.

Meg Yang: **So it's not just focusing in, say, general insurance or insurance specifically.**

Chris Dolman: That's right. I mean, that paper uses insurance as an example because, well, it's one that I understand quite well and I think the illustrations you can build are quite compelling from that, but the issues are certainly not restricted to insurance at all. We make that quite clear in the paper, but I do think insurance can be quite a good example to illustrate general principles sometimes. So yeah, that's what we decided to use.

Chris Dolman: What else can I talk about? Probably the other paper worth mentioning was last year's one that we did on our motor total loss project. So this is an example of a real-world case study that we published as IAG. We'd been part of the government's AI ethics principles pilot. And so, there's a shorter version of this on the department's website and a longer version that we published at the summit, which won the Melville Prize, which was quite nice to see. It's just a very accessible, practical case study, aiming to illustrate what people can actually do without resulting to rocket science or anything like that.

Chris Dolman: The methods used are quite basic, but it's really discussing how you practically build in ethical considerations, how you collaborate across the organisation to get something into market that can be quite valuable, because I think sometimes we get quite excited by nice, shiny things and new methods and what have you, which is great, but sometimes you can achieve quite a lot by doing something quite basic. So this really illustrates how that can be done and how you can apply ethics in a practical way in a real-world case study. So, we encourage more people to publish real-world case studies because they are quite helpful so that others can understand what companies are actually doing in practice.

Meg Yang: **Yeah, absolutely. I think those are interesting reads. So earlier you mentioned there's heaps of AI going wrong news in the news and we recently saw another one, which is about using facial recognition in retail shops, trying to capture the appearance of people that shops don't want in the shops and just keep them out in time. So from your view, what do you think can go wrong with adopting facial recognition algorithm and how to prevent or manage the risk?**

Chris Dolman: Yeah. Well, as you say, there's always stories of things going wrong as we've both said on this podcast, so this is another one that's made the news recently. But facial recognition has been in the news for quite a few years with people raising all sorts of concerns about it, so it's hardly a new topic. With this particular example, which is in retail, right? So the people were being filmed walking into shops and then facial recognition was being used to, I think, try and catch shoplifters or something like that, from what I read. So, the debate that I've seen that's ensued is really focused on privacy. It seems there's a debate now going on about whether that was just a breach of existing privacy regulations, nevermind whether or not that's something that we want to occur in general, but there's an allegation that potentially that might already be a problem under the rules as they currently are.

Chris Dolman: I think leaving aside the specifics of this, my general perspective on these sorts of discussions is that it really highlights that we do need some regulatory reform. That doesn't necessarily mean we need new regulation, although we might, and I know some people think we probably do with facial recognition, but I think we've failed when we get to this point, right? We're after the fact after something's happened having to make creative arguments on either side about whether something's fallen foul of existing rules or not, and we've failed at that point. What we need to do is be a lot more proactive.

Chris Dolman: So, I read white papers on facial recognition in retail settings probably five years ago, possibly more, where this sort of thing was being talked about. So I'm sure other people saw similar publications. At that point, we could have thought proactively about whether that was something we wanted and whether the existing rules that we've got clarified, whether that was acceptable or not. If the rules were unclear, we could have issued guidance to say one way or the other. If we're going to be proactive like that and I really think we'd need to be, rather than arguing after the fact, we'd need to resource regulators and industry appropriately to have discussions up front, but I do think there's promise in doing that because then you're deciding what you want to occur rather than trying to retrofit things after they've happened when damage has already been done, if it's damaging.

Meg Yang: **Absolutely. Yeah.**

Chris Dolman: I mean, if you're trying to encourage innovation, right? So one thing that slows down innovation is when the rules aren't clear because people then spend sometimes quite a lot of time trying to work that out. Whereas if you can be proactive, there's less hand-wringing then about whether your compliant or not, and you can just get on and innovate or not if it's not allowed and just not do it and not think about it, because it's not allowed and you know that. So it can really speed up the innovation cycle, which in that is something that plenty of people are keen to do in Australia. So

yeah, proactive regulation rather than reactive regulation is something I'm advocating for.

Meg Yang: **That's great. What do you think are the most topical issues, or in other words, what's hot in a space of ethical AI and responsible AI that you want to be proactive about?**

Chris Dolman: Yeah. Well, look, facial recognition, obviously, as we just talked about is a hot topic, so we need to work out what we're doing there. I think there's a broader related topic to that, which is so-called emotional AI, so trying to detect and interpret people's emotions or feelings by using... Especially video, people get very concerned about that, but also audio and text and what have you. So I'm sure you could take this podcast and try and work out whether I'm happy or sad or frustrated or angry or whatever. There are algorithms that we'll try and do this today.

Meg Yang: **That would be an interesting analysis.**

Chris Dolman: It certainly would be. I'm not convinced you'd get reliable results though, having looked at these things myself previously. So there's definitely accuracy issues with some of them. I mean, with faces, the idea that expressions on your face reliably predict your emotions is pretty dubious. There's been a lot of discussion on that the last few years. So it's a fairly dubious thing. One thing that I'm certainly aware of is that there's large cultural variation in how people communicate. So eye contact, for example, eye contact in some cultures is really encouraged when you're communicating and in others it's really frowned upon. So if your system is using eye contact as a gauge to work out whether someone's listening or communicating adequately, well, then it's not sensitive to culture and it's going to potentially lead to biases and all sorts of things going wrong.

Chris Dolman: So, that's a simple example of how these things can go wrong. There's obviously way more complex examples, but generally it's a fairly dubious idea that you can reliably and consistently across cultures predict emotions based on video or even voices and things like that. It just doesn't seem to work as well as perhaps it needs to for certain applications. For some applications, it probably doesn't matter too much, so it really depends on what you're doing, but certain application's failures are less acceptable than others. So, that's definitely one area which is being talked about a lot.

Chris Dolman: The other, and I mentioned this earlier, is data availability. So I mentioned Consumer Data Right, which is what we're doing in Australia, but just the general topic of data becoming more available is one that's being discussed a lot and whether or not that causes issues or not. We've just seen in the pandemic a real live example of that, right? So if the pandemic had happened 20 years ago, we wouldn't have had the data available that we had during COVID-19. Okay. But we had loads of data available to the

public, if you wanted to get access to it, to help you do your own analysis and make sense of what's going on, and loads of good came out of that, right? I saw all sorts of really interesting analysis and articles and apps getting built and all sorts of wonderful things came about, but it wasn't all good, right?

Chris Dolman: There were some negative things as well that arose from release of data. One of my favorite examples came from a particular world leader who said something along the lines of, "We need to stop testing people so then the numbers of infections goes down," which is just ludicrous, right? The measured number of infections might go down, of course, but not the actual number. So, there was an active incentive there to do something of general detriment because of the data that was being available because people didn't like the look of it. And so, we need to think quite carefully about what the release of data does in terms of incentives and things like that. It's not always going to be positive.

Chris Dolman: So, I'm going to be thinking about this topic a bit more and writing a paper for next year's International Congress on that sort of area, which should be of interest, hopefully. So then, yeah, what's hot? On a personal note,, I've been part of a successful Linkage Grant application, which is quite nice. So that's with colleagues from Australian National University, Sydney Uni and Gradient Institute. So we're going to kick off a three-year project researching responsible AI and data use in insurance, which is going to kick off pretty soon. So that's pretty exciting. I'll be sharing more on that over the next little while.

Meg Yang: **I'm looking forward to see that coming out. So with your experience and research and many others like yourself in working in the space of responsible and ethical AI, do you think the global environment for AI is improving over time and maybe people are gaining more clarity on how best to use AI responsibly? But if not, what do you think the future should look like?**

Chris Dolman: I think there's still way too much excitement, unfortunately, and maybe I'm just a party pooper, right? It's definitely true that there are more people raising concerns and risks now than there were five years ago, which is great, but I reckon there's even more people on the hype train as well. And so, is the balance right? Well, I'm not sure it is yet. I mean, in four or five years, since I've been watching the space actively, we have started to see better standards and regulation emerging, but it still feels like it's early days. It's been fairly slow and cumbersome.

Chris Dolman: We have seen what I consider to be quite a poor proposal out of the EU with its AI Act. I won't say in detail here, because I've said it a thousand times before, but I think it's not something that we should be copying for various reasons, which people can find out in my other work, I'm sure. I mean, rather than regulating mechanisms like AI, however defined, we

really just need to think about outcomes and whether those outcomes are appropriate and try and regulate in a neutral way, wherever possible, as much as we can do that.

Chris Dolman: So, as I said before, we just need to get proactive with our existing regulation. So, Australia has really, really good principles-based regulation in quite a lot of areas of the economy. First job is to make sure that those principles are well understood in the context of data availability, AI, automation, what have you. If we think we need new principles or we need to rejig them somewhat or clarify them or change them a bit, then we should get on and do that in a proactive way rather than sitting on our hands, and in five years time having an argument about something after harms have already been done. So, that's really what we need to start doing, I think, and then we can innovate in a way that's within the boundaries of what society desires. So that's what I would like to see, but I don't think we're yet there, unfortunately. So maybe we'll get there. Who knows?

Meg Yang: **So it sounds like there's still plenty of work that needs to be done. Why do you think actuaries should get involved and be aware of the issues?**

Chris Dolman: Well, at a very basic level, most actuaries use data. They might build models. They might influence decisions. They might affect customers. So whether or not you classify what you're doing as AI or big data or any of these other wonderful terms, the issues are still going to be there that I've talked about. So, as a professional, if you're using data and building models and affecting outcomes, you really should make yourself aware of this topic. It's something that is affecting your work. The users of your work are also going to have to be aware of it. So as a competent professional, you really need to be aware of some of these issues, at least at a high level, ideally in quite some detail.

Chris Dolman: Secondly, I just think actuaries generally take an interest in social issues. Quite a lot of actuaries I know have quite strong views on social issues and policy issues and the like, and I do think this topic is a great place for us to contribute and use our professional skills in a way that can benefit society. We have a quite unique combination of technical skills, risk management skills, knowledge about how businesses work and how decisions get made and our training and professional history on balancing competing interests can be brought to bear on these sorts of topics. So we have quite a lot of things we can bring to the topic. And so, I don't think we should be passive in the debate that's happening at the moment. We need to be quite active and we certainly can contribute a lot if we choose to be active. So I encourage others to be active.

Meg Yang: **Can you give some advice to other fellow actuaries who might be interested in this topic and maybe some resource they can refer to or how to contribute to any research in this area?**

Chris Dolman: Yeah. So I've spoken to heaps of people over the last few years in that sort of situation who were looking to get started. And so, if there's someone listening who wants to reach out and have a chat, I'm always happy to do that. If you want things to read, then, selfishly, I can say you can read the things that I've written, but you'll also find references to much better material within those papers as well, so maybe that gives you a place to start at least. There's plenty of guidance material that's been published by governments and agencies and regulators now, which is also a good place to start.

Chris Dolman: Some of the stuff out of Singapore has been very good recently, so maybe you have a look at that. The Veritas Project has some great documents that you can find, which gives you a guide on things like assessing fairness, so that's quite useful. I find some of the popular science to be quite good, particularly for folks like actuaries that are a bit more technically minded. You can find some popular science written by social scientists and the like, which will give you another angle to think through some of these issues. They're quite accessible usually. The ones that I find are most valuable are the ones that really focus on a couple of particular examples or issues to really pick apart what's going on and bring in the human side to what's going wrong. So, there's tons of examples now like that.

Chris Dolman: A good one I read recently was Automating Inequality, which if you need a place to start, it's probably a good one, but there's plenty of others as well. So that really helps technical people to open their minds to the impact of things that can go wrong. Look, if you want to contribute, well, Data Analytics Practice Committee is always looking for volunteers, so get in touch and we can find a way for you to contribute. I'm always looking for co-authors. So if there's something you want to write paper on, but you need someone to collaborate with, well, get in touch, let's have a chat and we'll see if we can do that.

Meg Yang: Thanks, Chris. Now, that's all we have time for today. Thanks for taking the time to join the podcast. Listeners, we hope you enjoy this discussion. You can listen to more Actuaries Institute podcast by visiting the podcast section of Actuaries Digital or via all mainstream podcast apps. Please write in with your questions and comments on the show. We'd love to hear from you. I'm Meg Yang. Bye for now.

© 2022 Actuaries Institute

Find out more about Actuaries and the Actuaries Institute:

- <https://www.actuaries.asn.au>

- <https://actuaries.digital>
- <https://linktr.ee/ActuariesInstitute>

Follow the Institute of Actuaries on our social channels:

- **LinkedIn:** <https://www.linkedin.com/company/792645/>
- **Facebook:** <https://www.facebook.com/pages/Actuaries-Institute/183337668450632>
- **Instagram:** <https://www.instagram.com/ActuariesInst>
- **Twitter:** <https://www.twitter.com/ActuariesInst>